

Accelerated Generalized Approximate Message Passing

Zilu Zhao, Fangqing Xiao, Jichao Chen, Dirk Slock,
Communication Systems Department, EURECOM, France
{zilu.zhao, fangqing.xiao, jichao.chen, dirk.slock}@eurecom.fr

Abstract—Generalized Approximate Message Passing (GAMP) enables Bayesian inference in linear models with non-identically and independently distributed (n.i.i.d.) priors and n.i.i.d. measurements of the linear mixture outputs. It represents an efficient technique for approximate inference, which becomes accurate when both rows and columns of the measurement matrix can be treated as sets of independent vectors and both dimensions become large. The fixed points of GAMP correspond to the extrema of a large system limit of the Bethe Free Energy (LSL-BFE), which represents a meaningful approximation criterion regardless of whether the measurement matrix exhibits the independence properties. However, the convergence of (G)AMP can be problematic for certain measurement matrices. In this paper, we revisit the LSL-BFE and its Lagrangian function. We derive an augmented GAMP algorithm by alternately enforcing the Karush-Kuhn-Tucker (KKT) conditions, called KKT-GAMP (KGAMP). To avoid matrix inversions, we introduce Adaptive (Accelerated) Gradient Descent (A(A)GD) techniques. Analysis shows convergence under relaxed conditions. Simulations indicate accelerated convergence compared to existing low complexity methods and illustrate the importance of adaptation.

I. INTRODUCTION

Low complexity Bayesian inference techniques are required for large dimensional sparse signal recovery. Sparse Bayesian Learning (SBL) is a Bayesian inference algorithm proposed by [1] and [2]. SBL is based on a hierarchical prior on the sparse signal, inducing sparsity in an underdetermined scenario by the fact of having to estimate the prior variance hyperparameters. The Linear Minimum Mean Squared Error (LMMSE) estimation step in SBL at each iteration involves matrix inversion, which makes it computationally complex [3]. The Approximate Message Passing (AMP) algorithm can be interpreted as Belief Propagation under the large system limit (LSL) [4], which represents messages by beliefs. Generalized AMP (GAMP) enables non-Gaussian priors and measurement processes. However, the convergence of (G)AMP can be problematic for some measurement/model matrices $\mathbf{A}$. Existing converging AMP versions before [5]: 1) adding the Alternating Direction Method of Multipliers (ADMM) [6] leading to a higher complexity ADMM-GAMP, 2) exploiting part of the singular value decomposition (SVD) of the measurement matrix in Vector AMP (VAMP) [7], [8] or esp. Unitarily Transformed UT-AMP [9] (but which do not allow to handle n.i.i.d. priors conveniently), 3) introducing damping [10], but with typically difficult to determine damping requirements.

- We propose a version of GAMP that alternately enforces the KKT conditions (KGAMP) of the large system limit of

the Bethe free Energy (LSL-BFE) [11].

- To avoid solving the least-squares subproblem through matrix inversion, we introduce Adaptive Accelerated Gradient Descent (AAGD).
- We provide a convergence analysis of the algorithm under relaxed conditions.
- Simulations show that the proposed KGAMP converges substantially faster than existing low-complexity methods, including AMBGAMP proposed in [5].

II. SYSTEM MODEL

We consider here a generalized linear model

$$p_{\mathbf{x}}(\mathbf{x}) = \prod_{i=1}^N p_{x_i}(x_i), \mathbf{z} = \mathbf{A}\mathbf{x}, p_{\mathbf{y}|\mathbf{z}}(\mathbf{z}) = \prod_{j=1}^M p_{y_j|z_j}(z_j), \quad (1)$$

where $\mathbf{A}$ is $M \times N$. We interpret the linear mixing as a conditional probability

$$p(\mathbf{z}|\mathbf{x}) = \delta(\mathbf{z} - \mathbf{A}\mathbf{x}). \quad (2)$$

In the signal recovery problem, we want to recover $\mathbf{x}$ from known $\mathbf{y}$ and $\mathbf{A}$. In later discussions, we introduce $\cdot$ as element-wise multiplication, $\cdot/$ as element-wise division, and $\mathbf{S} = \mathbf{A}.\mathbf{A}$.

III. LARGE SYSTEM LIMIT BETHE FREE ENERGY (LSL-BFE)

In [5], we arrived at a converging version of GAMP with a twist on ADMM, where, similarly to [10] and [6], we exploit a Lagrangian that is augmented with quadratic versions of the mean constraint while introducing an extra variable $\mathbf{u}$ for the estimated posterior mean of $\mathbf{x}$ to facilitate the alternating minimization. However, unlike ADMM for which the weights are almost arbitrary, the approach requires introducing very precise quadratic moment weights, which furthermore get ignored when optimizing over the corresponding variables. In the KKT approach pursued here, we can do away with the quadratic moment terms by introducing $\mathbf{u}$ in the variance expressions. Below, $\hat{\mathbf{x}} = \mathbb{E}[\mathbf{x}|\mathbf{q}_{\mathbf{x}}]$ and $\boldsymbol{\tau}_{\mathbf{x}} = \text{var}(\mathbf{x}|\mathbf{q}_{\mathbf{x}})$ will denote the vectors of means and variances of approximate posterior $\mathbf{q}_{\mathbf{x}}$, and similarly for $\mathbf{z}$.

After the LSL simplifications [15], the BFE with marginal pdf consistency constraints leading to belief propagation can

be seen to become equivalent to the following LSL-BFE with moment consistency constraints [5], [11]:

$$\begin{aligned} \min_{q_x, q_z, \tau_p} & D[q_x \| p_x] + D[q_z \| p_{y|z}] + H_G(q_z, \tau_p) \\ \text{s.t. } & \hat{\mathbf{z}} = \mathbf{A} \hat{\mathbf{x}}, \tau_p = \mathbf{S} \tau_x, \text{ where} \\ & H_G(q_z, \tau_p) = \frac{1}{2} \mathbf{1}^T (\tau_z \cdot / \tau_p) + \frac{1}{2} \sum_k \ln(\tau_{p_k}) \end{aligned} \quad (3)$$

Here $\mathbf{1}$ denotes a vector of 1's. We introduce the following mean square value (msv) definitions

$$\begin{aligned} (\mathbf{1} / \tau_r)^T \text{msv}(\mathbf{x} | q_x) &= \int \|\mathbf{x}\|_{\tau_r}^2 q_x(\mathbf{x}) d\mathbf{x} \\ (\mathbf{1} / \tau_p)^T \text{msv}(\mathbf{z} | q_z) &= \int \|\mathbf{z}\|_{\tau_p}^2 q_z(\mathbf{z}) d\mathbf{z} \end{aligned} \quad (4)$$

where $\|\mathbf{x}\|_{\tau}^2 = \mathbf{x}^T (\mathbf{x} / \tau)$ and $\text{msv}(\mathbf{x} | q_x)$ on a vector $\mathbf{x}$ operates elementwise. Also, $\text{var}(\mathbf{x} | q_x) = \text{msv}(\mathbf{x} - \mathbb{E} \mathbf{x} | q_x)$. The Lagrangian of (3) becomes

$$\begin{aligned} L &= D[q_x \| p_x] + D[q_z \| p_{y|z}] + \frac{1}{2} (\mathbf{1} / \tau_p)^T \text{msv}(\mathbf{z} - \mathbf{A} \mathbf{u} | q_z) \\ &+ \frac{1}{2} \sum_k \ln(\tau_{p_k}) + \mathbf{s}^T (\hat{\mathbf{z}} - \mathbf{A} \hat{\mathbf{x}}) - \frac{1}{2} \tau_s^T (\tau_p - \mathbf{S} \text{msv}(\mathbf{x} - \mathbf{u} | q_x)). \end{aligned} \quad (5)$$

Here, we have introduced a new optimization variable $\mathbf{u}$ to fluidify the alternating optimization, and dual variables $\mathbf{s}$ and τ_s . The proposed algorithm corresponds to alternating enforcing of the Karush–Kuhn–Tucker (KKT) conditions for (5). We define diagonal matrices $\mathbf{D}_p = \text{diag}(\tau_p)$, $\mathbf{D}_r = \text{diag}(\tau_r)$, and

$$\tau_r = \mathbf{1} / (\mathbf{S}^T \tau_s). \quad (6)$$

In (5), the weighting of $\text{msv}(\mathbf{x} - \mathbf{u} | q_x)$ can therefore be written as $\frac{1}{2} (\mathbf{1} / \tau_r)^T$, similar to the weighting of $\text{msv}(\mathbf{z} - \mathbf{A} \mathbf{u} | q_z)$. In the following discussion, we use $(\mathbf{D}_p, \mathbf{D}_r)$ and (τ_p, τ_r) interchangeably.

Proposition 0.1. *If the KKT conditions of (5) are satisfied, then $\forall \tau_p, \tau_r$*

$$\begin{aligned} (\mathbf{1} / \tau_p)^T \text{msv}(\mathbf{z} - \mathbf{A} \mathbf{u} | q_z) &= (\mathbf{1} / \tau_p)^T \text{var}(\mathbf{z} | q_z) \\ (\mathbf{1} / \tau_r)^T \text{msv}(\mathbf{x} - \mathbf{u} | q_x) &= (\mathbf{1} / \tau_r)^T \text{var}(\mathbf{x} | q_x). \end{aligned} \quad (7)$$

In the following, we use superscript $+$ to denotes solutions q_x^+ , q_z^+ , and $\mathbf{u}^+$ satisfying the KKT conditions. The main idea of the proof is to show that the KKT conditions of (5) imply

$$\mathbf{u}^+ = \mathbb{E}_{q_x^+}[\mathbf{x}]; \mathbf{A} \mathbf{u}^+ = \mathbb{E}_{q_z^+}[\mathbf{z}]. \quad (8)$$

Proof. We have in general

$$\begin{aligned} (\mathbf{1} / \tau_p)^T \text{msv}(\mathbf{z} - \mathbf{A} \mathbf{u} | q_z) &= (\mathbf{1} / \tau_p)^T \text{var}(\mathbf{z} | q_z) \\ &+ (\hat{\mathbf{z}} - \mathbf{A} \mathbf{u})^T \mathbf{D}_p^{-1} (\hat{\mathbf{z}} - \mathbf{A} \mathbf{u}), \end{aligned} \quad (9)$$

$$\begin{aligned} (\mathbf{1} / \tau_r)^T \text{msv}(\mathbf{x} - \mathbf{u} | q_x) &= (\mathbf{1} / \tau_r)^T \text{var}(\mathbf{x} | q_x) \\ &+ (\hat{\mathbf{x}} - \mathbf{u})^T \mathbf{D}_r^{-1} (\hat{\mathbf{x}} - \mathbf{u}). \end{aligned} \quad (10)$$

Due to the mean constraint, we have

$$\hat{\mathbf{z}}^+ = \mathbf{A} \hat{\mathbf{x}}^+, \text{ where } \hat{\mathbf{z}}^+ = \mathbb{E}_{q_z^+}[\mathbf{z}]; \hat{\mathbf{x}}^+ = \mathbb{E}_{q_x^+}[\mathbf{x}]. \quad (11)$$

As we can see, the quadratic terms in (9) and (10), and hence in (5) become zero and hence are minimized for $\mathbf{u}^+ = \hat{\mathbf{x}}^+$. Thus, at the optimal points, we have

$$\begin{aligned} \mathbb{E}_{q_z^+}[\|\mathbf{z} - \mathbf{A} \mathbf{u}^+\|_{\tau_p}^2] &= (\mathbf{1} / \tau_p)^T \text{var}(\mathbf{z} | q_z^+) \\ \mathbb{E}_{q_x^+}[\|\mathbf{x} - \mathbf{u}^+\|_{\tau_r}^2] &= (\mathbf{1} / \tau_r)^T \text{var}(\mathbf{x} | q_x^+). \end{aligned} \quad (12)$$

and hence (7). $\square$

To satisfy the KKT conditions, we not only set the partial derivatives of (5) w.r.t. $q_x(\mathbf{x})$, $q_z(\mathbf{z})$, τ_p , and $\mathbf{u}$ to zero, but also satisfy the equality constraints. The KKT conditions are equivalent to the following system of equations [5]

$$\frac{\partial \{ \|\hat{\mathbf{x}} - \mathbf{u}\|_{\tau_r}^2 + \|\hat{\mathbf{z}} - \mathbf{A} \mathbf{u}\|_{\tau_p}^2 \}}{\partial \mathbf{u}} = \mathbf{0} \quad (13)$$

$$\hat{\mathbf{z}}(\mathbf{s}) = \mathbf{A} \hat{\mathbf{x}}(\mathbf{s}) \quad (14)$$

$$q_x(\mathbf{x}) \propto p_x(\mathbf{x}) \mathcal{N}(\mathbf{x} | \mathbf{r}, \mathbf{D}_r) \quad (15)$$

$$\tau_p = \mathbf{S} \tau_x \quad (16)$$

$$q_z(\mathbf{z}) \propto p_{y|z}(\mathbf{z}) \mathcal{N}(\mathbf{z} | \mathbf{p}, \mathbf{D}_p) \quad (17)$$

$$\tau_s = (\mathbf{1} - \tau_z \cdot / \tau_p) \cdot / \tau_p \quad (18)$$

where

$$\begin{aligned} \hat{\mathbf{x}} &= \mathbb{E}_{q_x}[\mathbf{x}]; \hat{\mathbf{z}} = \mathbb{E}_{q_z}[\mathbf{z}]; \tau_x = \text{var}_{q_x}[\mathbf{x}]; \tau_z = \text{var}_{q_z}[\mathbf{z}] \\ \mathbf{p} &= \mathbf{A} \mathbf{u} - \mathbf{D}_p \mathbf{s}; \mathbf{r} = \mathbf{u} + \mathbf{D}_r \mathbf{A}^T \mathbf{s}; \tau_r = \mathbf{1} / (\mathbf{S}^T \tau_s). \end{aligned} \quad (19)$$

In (16) and (18), the variances τ_x , τ_z should in principle be the msv terms from (5), but we can replace them with the $\text{var}(\cdot)$ expressions due to the KKT conditions mentioned in the Proposition. This simplification corresponds to ignoring the dependence of the ADMM quadratic moment weights on the optimization variables τ_p , τ_s in [5], but at least now we know what this simplification corresponds to. In (14), we explicitly write $\hat{\mathbf{z}}(\mathbf{s})$, $\hat{\mathbf{x}}(\mathbf{s})$ because this is the constraint we use to update the Lagrange multiplier $\mathbf{s}$. Based on the above discussion, we will propose a low complexity algorithm in the following section.

IV. PROPOSED KGAMP ALGORITHM

The overall idea is to enforce the KKT conditions (13)-(18) alternately. We update 3 groups: $\mathbf{u}$, $(q_x, \tau_p, q_z, \mathbf{s})$, and τ_s in an alternating manner. We first update $\mathbf{u}$ based on (13). Then $(q_x, \tau_p, q_z, \mathbf{s})$ are updated by satisfying (14)-(17) at the same time. Finally, τ_s is updated based on (18).

A. Update of $\mathbf{u}$ by Adaptive Accelerated Gradient Descent

To avoid having to solve the linear system of equations in (13), consider at iteration t the quadratic cost function representing the terms in the Lagrangian (5) depending on $\mathbf{u}$:

$$F^t(\mathbf{u}) = \frac{1}{2} \|\hat{\mathbf{z}}^{t-1} - \mathbf{A} \mathbf{u}\|_{\tau_p^{t-1}}^2 + \frac{1}{2} \|\hat{\mathbf{x}}^{t-1} - \mathbf{u}\|_{\tau_r^{t-1}}^2. \quad (20)$$

We introduce a momentum auxiliary variable $\mathbf{w}$. The general form for (Nesterov) Accelerated Gradient Descent (AGD) is [16]

$$\begin{aligned} \mathbf{u}^t &= \mathbf{w}^{t-1} - \alpha \nabla F^t(\mathbf{w}^{t-1}) \\ \mathbf{w}^t &= \mathbf{u}^t + \beta (\mathbf{u}^t - \mathbf{u}^{t-1}), \end{aligned} \quad (21)$$

where the stepsize α and momentum gain β will be made adaptive (AAGD) by alternating line search. Define

$$\begin{aligned}\mathcal{H}^t &= \mathbf{D}(\mathbf{1}/\tau_r^{t-1}) + \mathbf{A}^T \mathbf{D}(\mathbf{1}/\tau_p^{t-1}) \mathbf{A} \\ \mathbf{g}^t(\mathbf{u}) &= \nabla F^t(\mathbf{u}) = \mathbf{g}^t(\mathbf{0}) + \mathcal{H}^t \mathbf{u} \\ &= -\mathbf{A}^T((\hat{\mathbf{z}}^{t-1} - \mathbf{A}\mathbf{u})/\tau_p^{t-1}) - (\hat{\mathbf{x}}^{t-1} - \mathbf{u})/\tau_r^{t-1}\end{aligned}\quad (22)$$

The quadratic cost function in (20) can be rewritten as

$$\begin{aligned}F^t(\mathbf{u}) &= \frac{1}{2} \mathbf{u}^T \mathcal{H}^t \mathbf{u} + \mathbf{u}^T \mathbf{g}^t(\mathbf{0}) + \text{const.} \\ \mathbf{g}^t(\mathbf{u}^{t-1}) &= \nabla F^t(\mathbf{u}^{t-1}) = \mathcal{H}^t \mathbf{u}^{t-1} + \mathbf{g}^t(\mathbf{0}).\end{aligned}\quad (23)$$

Due to the linearity of ∇F , (21) can be reformulated into a second-order recurrence for $\mathbf{u}$:

$$\begin{aligned}\mathbf{u}^t &= \mathbf{u}^{t-1} - \alpha (\mathcal{H}^t \mathbf{u}^{t-1} + \mathbf{g}^t(\mathbf{0})) + \beta (\mathbf{u}^{t-1} - \mathbf{u}^{t-2}) \\ &\quad + \alpha \beta \mathcal{H}^t (\mathbf{u}^{t-1} - \mathbf{u}^{t-2}).\end{aligned}\quad (24)$$

α and β should be adapted by minimizing $F^t(\mathbf{u}^t)$ after substituting (24) into $F^t(\mathbf{u})$ in (23). However, to avoid high-order equations, stepsize α and momentum gain β are optimized by alternating line search with a small inner loop. We have actually tried introducing an extra variable $\gamma = \alpha\beta$ and minimizing the resulting quadratic criterion in α , β , γ , but this does not appear to work as well.

1) *Adaptive Accelerated Gradient Descent Inner Loop:* We use t' to denote the inner loop iterations. Based on (24), define

$$\begin{aligned}\mathbf{d}_1^t(\beta) &= \mathbf{u}^{t-1} + \beta (\mathbf{u}^{t-1} - \mathbf{u}^{t-2}) \\ \mathbf{d}_2^t(\beta) &= (\mathcal{H}^t \mathbf{u}^{t-1} + \mathbf{g}^t(\mathbf{0})) + \beta \mathcal{H}^t (\mathbf{u}^{t-1} - \mathbf{u}^{t-2}) \\ \mathbf{d}_3^t(\alpha) &= \mathbf{u}^{t-1} - \alpha (\mathcal{H}^t \mathbf{u}^{t-1} + \mathbf{g}^t(\mathbf{0})) \\ \mathbf{d}_4^t(\alpha) &= (\mathbf{u}^{t-1} - \mathbf{u}^{t-2}) - \alpha \mathcal{H}^t (\mathbf{u}^{t-1} - \mathbf{u}^{t-2}).\end{aligned}\quad (25)$$

Therefore, (24) becomes

$$\mathbf{u}^t = \mathbf{d}_1^t(\beta) - \alpha \mathbf{d}_2^t(\beta) = \mathbf{d}_3^t(\alpha) + \beta \mathbf{d}_4^t(\alpha). \quad (26)$$

By setting the derivative of $F^t(\mathbf{u}^t|\alpha^{t'}, \beta^{t'-1})$ with respect to $\alpha^{t'}$ to zero, we obtain

$$\alpha^{t'} = \frac{\mathbf{d}_2^t(\beta^{t'-1})^T (\mathcal{H}^t \mathbf{d}_1^t(\beta^{t'-1}) + \mathbf{g}^t(\mathbf{0}))}{\mathbf{d}_2^t(\beta^{t'-1})^T \mathcal{H}^t \mathbf{d}_2^t(\beta^{t'-1})}. \quad (27)$$

Next, we set the derivative of $F^t(\mathbf{u}^t|\alpha^{t'}, \beta^{t'})$ w.r.t. $\beta^{t'}$ to zero, and we obtain

$$\beta^{t'} = -\frac{\mathbf{d}_4^t(\alpha^{t'})^T (\mathcal{H}^t \mathbf{d}_3^t(\alpha^{t'}) + \mathbf{g}^t(\mathbf{0}))}{\mathbf{d}_4^t(\alpha^{t'})^T \mathcal{H}^t \mathbf{d}_4^t(\alpha^{t'})}. \quad (28)$$

We propose to include an inner loop to find α and β iteratively based on (27) and (28). During the first sweep of the inner loop for AAGD, the complexity is $O(MN)$ due to the matrix-vector multiplication. However, after that, each consecutive sweep of the inner loop within one outer iteration has a complexity of $O(1)$.

B. Update of $q_x, \tau_p, q_z, \mathbf{s}$

By treating $\mathbf{u} = \mathbf{u}^t$ as given in (19), we have

$$\tau_p^t = \mathbf{S} \text{var}[\mathbf{x}|q_x^t]; \mathbf{p}^t = \mathbf{A}\mathbf{u}^t - \mathbf{D}_p^t \mathbf{s}^t; \mathbf{r}^t = \mathbf{u}^t + \mathbf{D}_r^{t-1} \mathbf{A}^T \mathbf{s}^t \quad (29)$$

Therefore, the updates for beliefs $q_x^t(\mathbf{x})$ and $q_z^t(\mathbf{z})$ in (15) and (17) are uniquely determined by the Lagrange multiplier $\mathbf{s}^t$. We find $\mathbf{s}^t$ by solving for $\mathbf{s}^t$ from the mean constraint

$$\mathbb{E}_{q_z^t|\mathbf{s}^t, \mathbf{u}^t}[\mathbf{z}] = \mathbf{A} \mathbb{E}_{q_x^t|\mathbf{s}^t, \mathbf{u}^t}[\mathbf{x}]. \quad (30)$$

However, solving (30) directly may lead to high complexity. We can also update $\mathbf{s}^t$ by line-search-aided gradient descent on a quadratic version of (30).

After that, we obtain the belief mean and variance by substituting $\mathbf{u}^t$ and $\mathbf{s}^t$ into q_x and q_z .

$$\begin{aligned}q_x^t(\mathbf{x}) &= p_x(\mathbf{x}) \mathcal{N}(\mathbf{x}|\mathbf{r}^t, \mathbf{D}_r^{t-1}) \\ q_z^t(\mathbf{z}) &= p_{y|\mathbf{z}}(\mathbf{z}) \mathcal{N}(\mathbf{z}|\mathbf{p}^t, \mathbf{D}_p^t)\end{aligned}\quad (31)$$

Therefore, we have

$$\begin{aligned}\hat{\mathbf{x}}^t &= \mathbb{E}_{q_x^t}[\mathbf{x}]; \tau_x^t = \text{var}_{q_x^t}[\mathbf{x}] \\ \hat{\mathbf{z}}^t &= \mathbb{E}_{q_z^t}[\mathbf{z}]; \tau_z^t = \text{var}_{q_z^t}[\mathbf{z}]\end{aligned}\quad (32)$$

1) *Updating $\mathbf{s}$ in the Gaussian Case:* Define diagonal matrices Σ_x and Σ_v of prior variances. With Gaussian p_x and $p_{y|\mathbf{z}}$

$$p_x(\mathbf{x}) = \mathcal{N}(\mathbf{x}|\mathbf{m}_x, \Sigma_x); p_{y|\mathbf{z}}(\mathbf{z}) = \mathcal{N}(\mathbf{z}|\mathbf{y}, \Sigma_v). \quad (33)$$

In the Gaussian case, the update of τ_p by (16) is independent of the value of $\mathbf{s}$. Therefore, we have

$$\begin{aligned}\tau_x^t &= \text{diag}[(\Sigma_x^{-1} + (\mathbf{D}_r^{t-1})^{-1})^{-1}] \\ \tau_p^t &= \mathbf{S} \tau_x^t; \mathbf{D}_p^t = \text{Diag}(\tau_p^t) \\ \tau_z^t &= \text{diag}[(\Sigma_v^{-1} + (\mathbf{D}_p^t)^{-1})^{-1}]\end{aligned}\quad (34)$$

Equation (30) is equivalent to

$$\mathbf{Q} \mathbf{s} = \mathbf{b}, \quad (35)$$

where

$$\begin{aligned}(\Sigma_p^t)^{-1} &= (\Sigma_v + \mathbf{D}_p^t)^{-1}; \mathbf{D}_z^t = \Sigma_v(\Sigma_v + \mathbf{D}_p^t)^{-1} \mathbf{D}_p^t \\ (\Sigma_r^t)^{-1} &= (\Sigma_x + \mathbf{D}_r^{t-1})^{-1}; \mathbf{D}_x^t = \mathbf{D}_r^{t-1}(\Sigma_x + \mathbf{D}_r^{t-1})^{-1} \Sigma_x \\ \mathbf{b}^t &= \mathbf{D}_p^t(\Sigma_p^t)^{-1} \mathbf{y} + \Sigma_v(\Sigma_p^t)^{-1} \mathbf{A} \mathbf{u}^t \\ &\quad - \mathbf{A} \mathbf{D}_r^{t-1}(\Sigma_r^t)^{-1} \mathbf{m}_x - \mathbf{A} \Sigma_x(\Sigma_r^t)^{-1} \mathbf{u}^t \\ \mathbf{Q}^t &= \mathbf{A} \mathbf{D}_x^t \mathbf{A}^T + \mathbf{D}_z^t.\end{aligned}\quad (36)$$

To solve (35) in a low complexity manner, we adopt line search (AGD) to the quadratic problem for which (35) can be considered the KKT condition

$$\arg \min_{\mathbf{s}} \frac{1}{2} \mathbf{s}^T \mathbf{Q}^t \mathbf{s} - \mathbf{s}^T \mathbf{b}^t. \quad (37)$$

Thus, by line search, $\mathbf{s}^t$ is updated by

$$\mathbf{s}^t = \mathbf{s}^{t-1} - \frac{(\mathbf{h}^t)^T \mathbf{h}^t}{(\mathbf{h}^t)^T \mathbf{Q}^t \mathbf{h}^t} \mathbf{h}^t \quad (38)$$

where $\mathbf{h}^t = \mathbf{Q}^t \mathbf{s}^{t-1} - \mathbf{b}^t$.

2) *Updating s in the General Case:* In this case we can construct equivalent Gaussian priors by the Expectation Propagation rule, by dividing the Gaussian posterior models constructed from the MMSE estimate and associated MSE, by the Gaussian extrinsics:

$$\tilde{p}_x(\mathbf{x}) = \mathcal{N}(\mathbf{x}|\mathbf{m}_x, \Sigma_x) \propto \frac{\mathcal{N}(\mathbf{x}|\hat{\mathbf{x}}, \mathbf{D}_x)}{\mathcal{N}(\mathbf{x}|\mathbf{r}, \mathbf{D}_r)} \quad (39)$$

$$\tilde{p}_z(\mathbf{z}) = \mathcal{N}(\mathbf{z}|\mathbf{m}_z, \Sigma_z) \propto \frac{\mathcal{N}(\mathbf{z}|\hat{\mathbf{z}}, \mathbf{D}_z)}{\mathcal{N}(\mathbf{z}|\mathbf{p}, \mathbf{D}_p)} \quad (40)$$

C. Update of τ_s

Finally, we update τ_s from $\partial L / \partial \tau_p = 0$, leading to (18):

$$\tau_s^t = (1 - \tau_z^t / \tau_p^t) / \tau_p^t. \quad (41)$$

By definition in (19), $\tau_r^t = \mathbf{1} / (\mathbf{S}^T \tau_s^t)$, which actually follows from asymptotic Belief Propagation [15].

V. CONVERGENCE ANALYSIS

We consider a relaxed algorithm in which,

- Both p_x and $p_{y|z}$ are Gaussian, as in (33),
- $\mathbf{u}^t$ is obtained by zeroing the gradient of (20):

$$\mathbf{u}^t = (\mathcal{H}^t)^{-1} \mathbf{g}(\mathbf{0}), \quad (42)$$

- $\mathbf{s}^t$ is obtained by solving (30), and hence (35).

It has been shown in [5] that the variance subsystem (which is independent of the mean subsystem in the Gaussian prior case) is guaranteed to converge. In [19], it is shown that under the large system assumption, τ_x^∞ converges to the LMMSE posterior variances:

$$\tau_x^\infty = \text{diag}[(\mathbf{A}^T \Sigma_v^{-1} \mathbf{A} + \Sigma_x^{-1})^{-1}]. \quad (43)$$

In the following, we will show that the mean subsystem of the relaxed version of KGAMP also converges.

Although this is not essential for the convergence analysis, as the variance subsystem converges fairly fast, let's assume that it has converged already: $\forall t, \tau_p^t = \tau_p$, and $\tau_r^t = \tau_r$, and hence also $\mathbf{Q}^t := \mathbf{Q}$. At the end of each iteration, due to the second and third relaxation assumptions, and (30), we have

$$\hat{\mathbf{z}}^t = \mathbf{A} \hat{\mathbf{x}}^t; \mathbf{u}^t = \hat{\mathbf{x}}^{t-1}. \quad (44)$$

Furthermore, with the second relaxation assumption, we have

$$\mathbf{s}^t = \mathbf{Q}^{-1} \mathbf{b}^t, \quad (45)$$

On the other hand, with fixed τ_p , τ_r , and the Gaussian prior, we have

$$\hat{\mathbf{x}}^t = \mathbf{D}_r(\Sigma_x + \mathbf{D}_r)^{-1} \mathbf{m}_x + \Sigma_x(\Sigma_x + \mathbf{D}_r)^{-1} \mathbf{r}^t, \quad (46)$$

where $\mathbf{r}^t$ is computed in (29). Substituting (44) and (45) into (29) and substituting the resulting (29) into (46), we have

$$\hat{\mathbf{x}}^t = \Phi \hat{\mathbf{x}}^{t-1} + \text{const}, \quad (47)$$

where

$$\begin{aligned} \Phi &= (\mathbf{A}^T \mathbf{D}_z^{-1} \mathbf{A} + \mathbf{D}_x^{-1})^{-1} (\mathbf{A}^T \mathbf{D}_p^{-1} \mathbf{A} + \mathbf{D}_z^{-1}), \\ \mathbf{D}_x &= (\mathbf{D}_r^{-1} + \Sigma_x^{-1})^{-1}; \\ \mathbf{D}_z &= (\mathbf{D}_p^{-1} + \Sigma_v^{-1})^{-1}. \end{aligned} \quad (48)$$

We use $\Theta_1 \leq \Theta_2$ to denote that $\Theta_2 - \Theta_1$ is positive semi-definite. From (48), we can verify that $\mathbf{D}_z \leq \mathbf{D}_p$ and $\mathbf{D}_x \leq \mathbf{D}_r$. Therefore, we have $\Phi \leq \mathbf{I}$. Hence, the update from $\hat{\mathbf{x}}^{t-1}$ to $\hat{\mathbf{x}}^t$ is a contraction and converges.

VI. SIMULATIONS

A. Gaussian Case

We consider an $M \times N = 5 \times 10$ measurement model. The input signal follows an exponential decay $x_n \sim \mathcal{N}(0, \sigma_{x_n}^2)$, where $\sigma_{x_n}^2 = 0.01^{(n-1)}$. The elements of $\mathbf{A}$ are drawn i.i.d. $\mathcal{N}(0, \frac{1}{N})$. We consider white Gaussian noise with power $\sigma_v^2 = 10^{-4}$. We randomly generate 1000 realizations of $\mathbf{x}$, $\mathbf{v}$ and $\mathbf{A}$. For AAGD, we adopt an inner loop of 50 iterations. Since the novelty of KGAMP (w.r.t, AMBGAMP) is in the mean subsystem and we limit the simulations to the Gaussian case, we first obtain τ_r and τ_p by running the variance updates for 500 iterations. After that, we keep them fixed during the mean subsystem iterations. Furthermore, the same set of τ_r and τ_p is used in different methods to be compared.

In addition to the benchmark methods ADMM [20] and AMBGAMP [5], we also compare different low complexity methods for updating $\mathbf{u}$ in (31).

In Fig. 1, we plot the normalized mean squared error (NMSE) with respect to the ground-truth $\mathbf{x}$. Whereas in Fig. 2, we plot the NMSE to the MMSE estimate

$$\hat{\mathbf{x}}_{MMSE} = (\mathbf{A}^T \Sigma_v^{-1} \mathbf{A} + \Sigma_x^{-1})^{-1} \mathbf{A}^T \Sigma_v^{-1} \mathbf{y}. \quad (49)$$

In all “KGAMP-” based algorithms, we only update $\mathbf{u}$ once per iteration. The proposed method is “KGAMP” in the figures, short for “KGAMP-AAGD”. In “KGAMP-Nesterov” and “KGAMP-GD”, we use a fixed stepsize $1/L$, where $L = \max[\text{eig}(\mathbf{Q}^t)]$. “KGAMP-Nesterov” uses standard Nesterov accelerated Gradient descent for updating $\mathbf{u}$, while “KGAMP-GD” uses the standard gradient descent algorithm for updating $\mathbf{u}$. “KGAMP-GD-line-search” (which is “KGAMP-AGD” in fact) uses gradient descent with line search stepsize for updating $\mathbf{u}$ [5]. The dashed curves are methods that require matrix inversion. “KGAMP Least Squares” uses (45) to update $\mathbf{s}$ and hence $\mathbf{u}^t = \hat{\mathbf{x}}^{t-1}$ to update $\mathbf{u}$.

“KGAMP Least Squares” provides an indication of how fast the algorithm can converge, as only the alternating optimization component remains to be further accelerated. ADMM-GAMP converges at a comparable rate; however, more realistic gradient-based implementations exhibit slower convergence than the other algorithms. The simulations also demonstrate that optimizing the stepsizes in Nesterov's method leads to substantial additional acceleration.

B. Gaussian Mixture Model Prior

We consider a $\mathbf{A} \in \mathbb{R}^{5 \times 10}$ system. In this case the input elements are independent and follow a decaying Gaussian mixture model (GMM) distribution:

$$\begin{aligned} x_n &= w_{n,1} \mathcal{N}(0.5^n m_{n,1}, [0.5 \sigma_{n,1}]^{2n}) \\ &\quad + w_{n,2} \mathcal{N}(0.5^n m_{n,2}, [0.5 \sigma_{n,2}]^{2n}), \end{aligned} \quad (50)$$

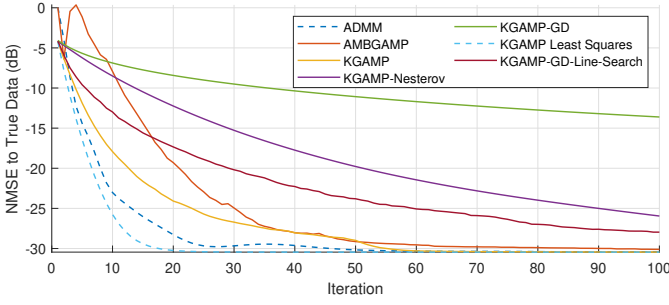


Fig. 1. NMSE-to- $\mathbf{x}$ vs Iterations

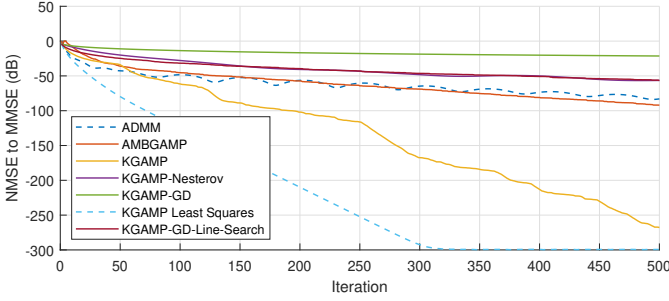


Fig. 2. NMSE-to- $\hat{\mathbf{x}}_{MMSE}$ vs Iterations

where $\forall n, i$, the $w_{n,i}$ are drawn from a uniform distribution between 0 and 1. The sum $w_{n,1} + w_{n,2}$ gets normalized after drawing the realizations. The Gaussian parameters $m_{n,i}$, $\sigma_{n,i}$ are drawn from a standard Gaussian distribution. All the parameters are known during the signal recovery process. The measurement matrix $\mathbf{A}$ follows the same model as in the Gaussian case above, and the noise is again white Gaussian but with variance adjusted to attain SNR=17dB.

We perform 400 different realizations with different prior distribution, measurement matrix and noise. To evaluate the accuracy, we compare the normalized mean squared error to MMSE estimates (which can be computed analytically for a GMM).

In Fig. 3 we compare the proposed KGAMP algorithm to the ADMM-GAMP algorithm proposed in [6], in which the order of some variance updates (involving τ_p) needs to be modified (as in [5]) for the algorithm to become convergent in our simulations, and this with a single loop algorithm (unlike the double loop in [6]). In both algorithms, $\mathbf{u}$ gets updated by solving the least-squares problem. The simulation shows that KGAMP converges faster and appears more stable after convergence.

VII. CONCLUSIONS

We propose KGAMP based on alternately enforcing the KKT conditions of the LSL-BFE. To avoid solving the systems of linear equations through matrix inversion, we introduce Adaptive (Accelerated) Gradient Descent (A(A)GD). With Gaussian or non-Gaussian priors and measurements, the proposed algorithm converges (very fast) when the KKT conditions in each alternating step are satisfied exactly. In low

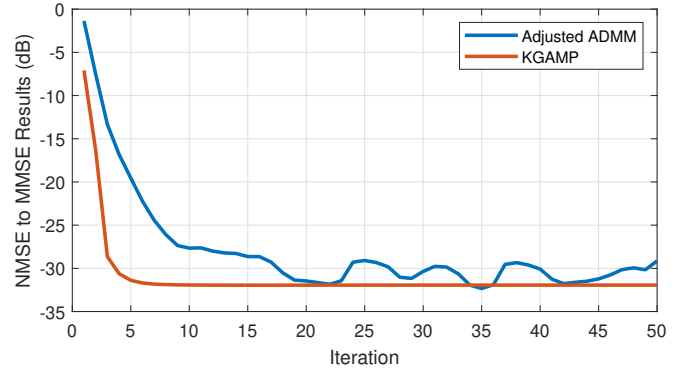


Fig. 3. NMSE-to- $\hat{\mathbf{x}}_{MMSE}$ vs Iterations

complexity GAMP versions, convergence can be slowed down due to (i) GD for updating $\mathbf{u}$, (ii) (GD) for updating $\mathbf{s}$, (iii) or a GD-style update for $\mathbf{s}$ in ADMM algorithm versions. It appears that (i) is the dominating slow down factor. Finally, there is also the alternating optimization nature of GAMP which may pose the ultimate convergence speed challenge.

Proposed more classical approach to

Before pursuing the approach presented in this paper, we actually tried accelerating the convergence of the AMBGAMP algorithm in [5] by pursuing accelerated ADMM algorithms, which have been introduced over the last decade. However, in spite of the Lagrange multiplier update in ADMM corresponding to a gradient ascent step, the stepsize has to be chosen very precisely for ADMM. This makes any deviation in order to introduce acceleration quite delicate. On the other hand, ADMM is used very often in distributed algorithms in which local versions of the variables are updated, among which then equality is imposed via ADMM. The more classical Lagrangian approach proposed in this paper may open new avenues for developing distributed algorithms with accelerated convergence.

ACKNOWLEDGEMENTS

EURECOM's research is partially supported by its industrial members: ORANGE, BMW, SAP, iABG, Norton LifeLock, by the French PEPR-5G projects PERSEUS and YACARI, the EU H2030 project CONVERGE, and by a Huawei France funded Chair towards Future Wireless Networks.

REFERENCES

- [1] M. E. Tipping, "Sparse Bayesian Learning and the Relevance Vector Machine," *J. Mach. Learn. Res.*, vol. 1, 2001.
- [2] D. P. Wipf and B. D. Rao, "Sparse Bayesian Learning for Basis Selection," *IEEE Trans. on Sig. Proc.*, vol. 52, no. 8, Aug. 2004.
- [3] C. K. Thomas and D. Slock, "Low Complexity Static and Dynamic Sparse Bayesian Learning Combining BP, VB and EP Message Passing," in *Asilomar Conf. on Sig., Sys., and Comp.*, CA, USA, 2019.
- [4] Z. Zhao and D. Slock, "Extrinsics and Linearized Component-Wise Conditionally Unbiased MMSE Estimation in Approximate Message Passing," in *Int'l Conf. Computing, Networking and Communications (ICNC)*, 2025.

- [5] C. K. Thomas, Z. Zhao, and D. Slock, "Towards Convergent Approximate Message Passing by Alternating Constrained Minimization of Bethe Free Energy," in *IEEE Information Theory Workshop (ITW)*, 2023.
- [6] S. Rangan, A. Fletcher, P. Schniter, and U. Kamilov, "Inference for Generalized Linear Models via Alternating Directions and Bethe Free Energy Minimization," *IEEE Trans. Info. Theory*, Jan. 2017.
- [7] S. Rangan, P. Schniter, and A. K. Fletcher, "Vector Approximate Message Passing," *IEEE Trans. On Info. Theo.*, vol. 65, no. 10, Oct. 2019.
- [8] P. Schniter, S. Rangan, and A. K. Fletcher, "Vector approximate Message Passing for the Generalized Linear Model," in *In IEEE 50th Asilomar Conference on Signals, Systems and Computers*, 2016.
- [9] Q. Guo and J. Xi, "Approximate Message Passing with Unitary Transformation," 2015. [Online]. Available: <https://arxiv.org/abs/1504.04799>
- [10] S. Rangan, P. Schniter, A. Fletcher, and V. Cevher, "Fixed Points of Generalized Approximate Message Passing with Arbitrary Matrices," *IEEE Trans. Info. Theory*, Dec. 2016.
- [11] Z. Zhao and D. Slock, "Bethe Free Energy and Extrinsic in Approximate Message Passing," in *Asilomar Conf. Signals, Systems, and Computers*, 2023.
- [12] Z. Zhao, F. Xiao, and D. Slock, "Approximate Message Passing for Not So Large niid Generalized Linear Models," in *IEEE Int'l Workshop on Signal Processing Advances in Wireless Communications (SPAWC)*, 2023.
- [13] Z. Zhao, F. Xiao, C. K. Thomas, and D. Slock, "Reconciling AMP Algorithms derived from Belief Propagation or the Large System Limit Bethe Free Energy," in *IEEE Int'l Conf. Acoustics, Speech and Signal Processing (ICASSP)*, 2025.
- [14] Z. Zhao, F. Xiao, and D. Slock, "Vector Approximate Message Passing for Not So Large N.I.I.D. Generalized I/O Linear Models," in *IEEE Int'l Conf. Acoustics, Speech and Signal Processing (ICASSP)*, 2024.
- [15] —, "Extrinsic and Linearized Component-Wise Conditionally Unbiased MMSE Estimation as in GAMP," in *Asilomar Conf. Signals, Systems, and Computers*, 2024.
- [16] —, "Accelerated Alternating Optimization of MIMO Transceivers," in *Asilomar Conf. Signals, Systems, and Computers*, 2025.
- [17] —, "Expectations in Expectation Propagation," 2026. [Online]. Available: <https://arxiv.org/abs/2512.08034>
- [18] Z. Zhao and D. Slock, "Approximating Univariate Factored Distributions via Message-Passing Algorithms," 2026. [Online]. Available: <https://arxiv.org/abs/2602.01377>
- [19] —, "Variance Predictions in VAMP/UAMP with Right Rotationally Invariant Measurement Matrices for niid Generalized Linear Models," in *European Signal Processing Conf. (EUSIPCO)*, 2023.
- [20] S. Rangan, A. K. Fletcher, P. Schniter, and U. S. Kamilov, "Inference for Generalized Linear Models via Alternating Directions and Bethe Free Energy Minimization," *IEEE Transactions on Information Theory*, vol. 63, no. 1, pp. 676–697, 2017.