

PROVABLY CONVERGENT NONCONVEX ALGORITHM FOR VOLUME OPTIMIZATION-BASED LATENT COMPONENT ANALYSES

SUPPLEMENTAL MATERIAL

Anonymous authors

Paper under double-blind review

A RELATED WORKS

In this section we introduce some existing algorithms that directly tackles the formulation, repeated here:

$$\underset{\mathbf{W}}{\text{minimize}} \quad -\frac{1}{2} \log \det \mathbf{W}\mathbf{W}^\top + g(\mathbf{W}\mathbf{X}). \quad (1)$$

More specifically, we introduce two algorithms that rely on existing off-the-shelf toolboxes to solve linear programming problems as a sub-routine in each iteration. There exists an early work based on augmented Lagrangian method specifically designed for ACA (Bioucas-Dias, 2009). We will include it in the experiment comparing all these algorithms, but the description is omitted as it is somewhat similar in the spirit of our proposed L-ADMM.

A.1 BLOCK COORDINATE DESCENT (BCD)

This method only works when $\mathbf{W}$ is square, so $\det \mathbf{W}\mathbf{W}^\top = |\det \mathbf{W}|^2$. Using the Laplace’s formula, we know that $\det \mathbf{W}$ is a linear function with respect to the i th row of $\mathbf{W}$ using the co-factor expansion

$$\det \mathbf{W} = \sum_{j=1}^k (-1)^{i+j} W_{ij} \det \mathbf{W}_{ij},$$

where the matrix $\mathbf{W}_{ij}$ is obtained by deleting the i th row and j th column of $\mathbf{W}$. Then to apply the block coordinate descent algorithm (Bertsekas, 1999), one would minimize $-|\mathbf{f}_i^\top \mathbf{w}_i| + g_i(\mathbf{w}_i^\top \mathbf{X})$, where $\mathbf{f}_i \in \mathbb{R}^k$ is a vector defined according to the co-factor expansion formula, and g_i is a function in $\mathbb{R}^k$ by fixing all except the i th row of $\mathbf{W}$.

It looks as if we need to solve two linear programs for each row update, one to maximize $\mathbf{f}_i^\top \mathbf{w}_i$ and one to minimize $\mathbf{f}_i^\top \mathbf{w}_i$, but it turns out one only needs to maximize $\mathbf{f}_i^\top \mathbf{w}_i$. For BCA and SCA where there is a sign ambiguity in each component, either solution is a sign flip of the other and equally good; for NCA and ACA, due to the nonnegativity constraint, the sign of $\mathbf{f}_i^\top \mathbf{w}_i$ should equal to that of $\det \mathbf{W}$. Furthermore, recall that the Cramer’s rule shows that

$$[\mathbf{P}^{-1}]_{ji} = (-1)^{i+j} \det \mathbf{P}_{ij} / \det \mathbf{P},$$

meaning we can simply define $\mathbf{f}_i$ as the i th column of $\mathbf{P}^{-1}$, and it would not affect the row updates.

The idea was first proposed to solve ACA by Chan et al. (2009). Several follow-up works have been proposed to solve related problems, such as NCA (Huang et al., 2016), ACA (Huang & Fu, 2019), BCA (Hu & Huang, 2023b), and SCA (Hu & Huang, 2023a). Notice that BCD works the best when the constraints are all separable over the blocks (Bertsekas, 1999); this makes BCD works fairly good on BCA, SCA, and NCA, but not particularly well on ACA (even though it was first proposed on this problem).

A.2 FRANK-WOLFE (FW)

The Frank-Wolfe algorithm, also known as the conditional gradient method for constrained optimization (Bertsekas, 1999), iteratively minimizes a linear objective, defined by the gradient at the

Algorithm 1 Solving (1) with BCD

```

initialize  $\mathbf{W}_{(0)}$ 
repeat
  for  $i = 1, \dots, k$  do
     $\mathbf{f} = \mathbf{W}^{-1} \mathbf{e}_i$ 
     $\mathbf{w} = \arg \min_{\mathbf{w}} -\mathbf{f}^\top \mathbf{w} + g_i(\mathbf{w}_i^\top \mathbf{X})$ 
    replace  $i$ th row of  $\mathbf{W}$  with  $\mathbf{w}^\top$ 
  end for
until convergence

```

current iterate, under the same constraint set to determine the search direction and obtain the next iterate via some line search approach along the search direction. For the log-determinant objective in (1), we have that the gradient is $-(\mathbf{P}^\dagger)^\top$. As a result, the Frank-Wolfe algorithm for (1) is given in Algorithm 2.

Algorithm 2 Solving (1) with Frank-Wolfe

```

initialize  $\mathbf{P}_{(0)}$ 
for  $t = 0, 1, 2, \dots$  until convergence do
   $\mathbf{W}_d = \arg \min_{\mathbf{W}} -\text{Tr}(\mathbf{W}_{(t)}^\dagger \mathbf{W}) + g(\mathbf{W}\mathbf{X})$ 
   $\alpha \leftarrow 1$ 
  while  $-\log |\det(\mathbf{W}_{(t)} + \alpha(\mathbf{W}_d - \mathbf{W}_{(t)}))| > -\log |\det \mathbf{W}_{(t)}| + (\alpha/2) \text{Tr}(\mathbf{W}_{(t)}^{-1}(\mathbf{W}_d - \mathbf{W}_{(t)}))$  do
     $\alpha \leftarrow \alpha/2$ 
  end while
   $\mathbf{W}_{(t+1)} = \mathbf{W}_{(t)} + \alpha(\mathbf{W}_d - \mathbf{W}_{(t)})$ 
end for

```

Regarding the line search step, the backtracking line search (Armijo rule) (Bertsekas, 1999) is used to guarantee sufficient decrease of the objective function. Since g is convex, as long as $\mathbf{W}_{(t)}$ is feasible, then $\mathbf{W}_{(t+1)}$ is also feasible since it is a convex combination of $\mathbf{W}_{(t)}$ and $\mathbf{P}_d$, which are by definition feasible. Therefore, the only nontrivial part is to find a feasible initialization $\mathbf{P}_{(0)}$. This can be done by optimizing an arbitrary linear objective subject to g (which means one should not apply line search at this step).

Frank-Wolfe was proposed to solve an instance of (1) by Hu & Huang on SCA (2023a) and NCA/ACA (2024).

The two proposed algorithms 2 and 1 shows striking similarities, especially considering they both fundamentally solve linear programs in each iteration. For nonconvex optimization, both algorithms guarantees that every limit point is a stationary point with some additional assumptions, which is reassuring to know. The differences are as follows:

1. Block coordinate descent is guaranteed to monotonically improve the objective function by design, so there is no need for line search as in Frank-Wolfe. This could save some computation as calculating $\det \mathbf{W}$ may not be cheap when k is large.
2. On the other hand, each iteration of Frank-Wolfe only need to calculate $\mathbf{W}^\dagger$ once for the update of all its k rows, whereas BCD requires to recalculate $\mathbf{W}^{-1}$ for each of its row updates. Overall, the per-iteration complexity is almost identical.

A.3 NUMERICAL PERFORMANCES

We now provide some numerical experiments to demonstrate how our proposed L-ADMM performs compared to the existing algorithms outlined in this section. The settings are exactly the same as in §4, except that in the figures the horizontal axis shows time elapsed in seconds, not iterations—while L-ADMM takes hundreds to thousands of iterations to converge, neither BCD nor FW takes more than a few tens, but their per-iteration complexities are much higher as they require solving linear programs as a sub-routine.

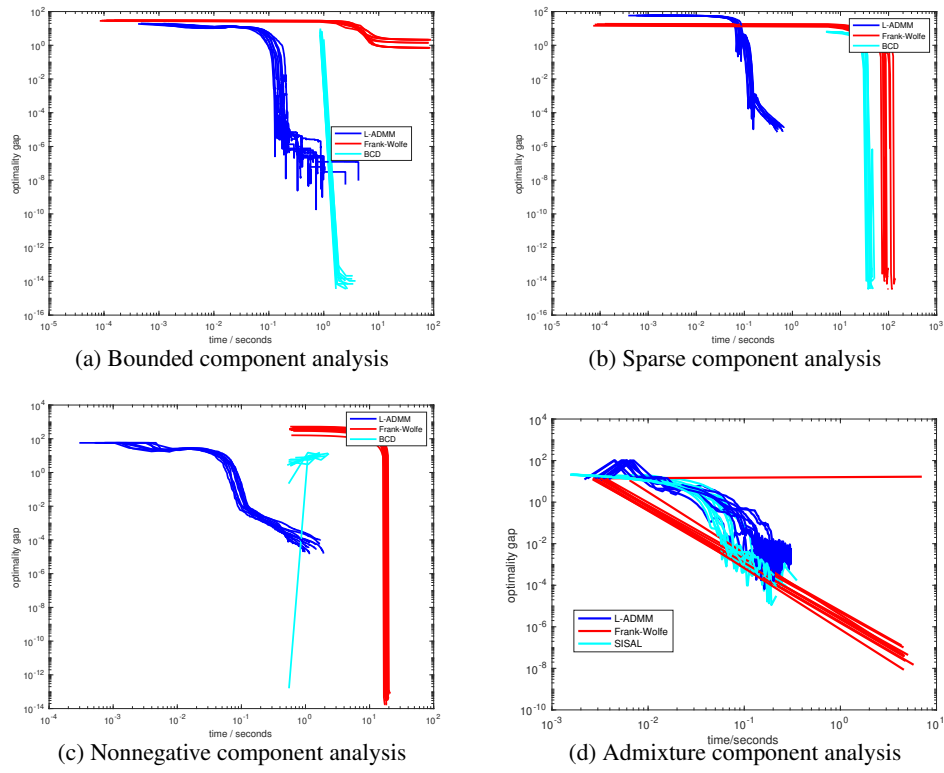


Figure 1: Convergence of L-ADMM vs. BCD and FW for various latent component analyses models on 10 random trials.

The convergence behavior of 10 random trials of L-ADMM, BCD, and FW are shown for BCA, SCA, and NCA in Figure 1. For ACA, we find that BCD proposed by Chan et al. (2009) takes too long to converge in all cases due to the fact that the constraints are not separable over the rows, so BCD is not included; instead, we include another augmented Lagrangian method called SISAL designed specifically for ACA as comparison. As we can see, L-ADMM works significantly better than BCD and FW in all cases. In some cases, BCD or FW may not even solve the problem, such as FW for BCA or BCD for NCA, which is in fact reasonable as we are trying to solve a nonconvex problem—what is incredible is that L-ADMM always works within a very short amount of time, which motivates this work. L-ADMM is overall comparable to the highly efficient SISAL, which is a dedicated algorithm for ACA, while L-ADMM is also highly flexible for other problems.

B PROOF OF LEMMA 3.2

Recall the proposed L-ADMM is

$$\begin{cases} \mathbf{W}_{t+1} \leftarrow (\mathbf{S}_t + \mathbf{U}_t + \gamma \mathbf{S}_t^{\dagger\top}) \mathbf{X}^\dagger \\ \mathbf{S}_{t+1} \leftarrow \text{Prox}_{\gamma g}(\mathbf{W}_{t+1} \mathbf{X} - \mathbf{U}_t) \\ \mathbf{U}_{t+1} \leftarrow \mathbf{U}_t + \mathbf{S}_{t+1} - \mathbf{W}_{t+1} \mathbf{X} \end{cases} \quad (2)$$

and Lemma 3.2 is

Lemma B.1. *When running the L-ADMM iterations, if $\mathbf{S}_t$ satisfies*

$$\log |\det \mathbf{S}_t \mathbf{S}_t^\dagger| \leq \text{Tr} \mathbf{S}_t \mathbf{S}_t^\dagger - k, \quad (3)$$

then $\mathbf{S}_{t+1}$ also satisfies

$$\log |\det \mathbf{S}_{t+1} \mathbf{S}_{t+1}^\dagger| \leq \text{Tr} \mathbf{S}_{t+1} \mathbf{S}_{t+1}^\dagger - k,$$

It is well known that $\text{Tr}(\mathbf{S}_t \mathbf{S}_t^\dagger)$ equals to the sum of its eigenvalues, and $\det(\mathbf{S}_t \mathbf{S}_t^\dagger)$ equals to the product of its eigenvalues. Since we do not restrict $\mathbf{S}_t \mathbf{S}_t^\dagger$ to be symmetric, its eigenvalues may be complex. However, we do restrict $\mathbf{S}_t \mathbf{S}_t^\dagger$ to be real, so its complex eigenvalues always come in conjugate pairs. This means if a complex eigenvalue takes the form $\rho + i\eta$, then $\rho - i\eta$ is also an eigenvalue.

A sufficient condition for (3) is that all eigenvalues satisfy

$$\rho - 1 \geq \frac{1}{2} \log(\rho^2 + \eta^2). \quad (4)$$

This inequality obviously holds if $\eta = 0$ (i.e., this eigenvalue is real) when $\rho > 0$ as per the famous inequality $\rho - 1 \geq \log \rho$. What is somewhat less obvious is that this inequality also holds when $\rho > |\eta|$ when $\eta \neq 0$. It is easy to verify that for a fixed η , the second derivative of $-(1/2) \log(\rho^2 + \eta^2)$ with respect to ρ is positive when $\rho > |\eta|$. Thus invoking the first-order condition for a convex function at point 1, we obtain the inequality (4).

The resulting condition $\Re(\lambda) > |\Im(\lambda)|$ for every eigenvalue λ of the matrix $\mathbf{S}_t \mathbf{S}_t^\dagger$ defines a region on the complex plane in which it can reside in order for (3) to hold. This region is illustrated in Figure 2. Note that if (4) is true for $\mathbf{S}_t \mathbf{S}_t^\dagger$, it is trivially true for its inverse $(\mathbf{S}_t \mathbf{S}_t^\dagger)^{-1}$ as well.

Proof of Lemma 3.2. Replacing $\mathbf{X} = \mathbf{A}_t \mathbf{S}_t$ in (2), we first obtain from the $\mathbf{W}$ update that

$$\mathbf{W}_{t+1} \mathbf{A}_t = (\mathbf{S}_t + \mathbf{U}_t) \mathbf{S}_t^\dagger + \gamma (\mathbf{S}_t \mathbf{S}_t^\top)^{-\top},$$

and the $\mathbf{U}$ update gives

$$\mathbf{U}_{t+1} \mathbf{S}_t^\dagger = (\mathbf{U}_t + \mathbf{S}_{t+1}) \mathbf{S}_t^\dagger - \mathbf{W}_{t+1} \mathbf{A}_t.$$

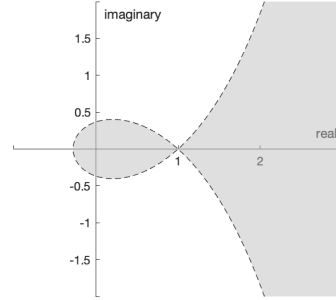


Figure 2: If all the eigenvalues of $\mathbf{S}_t \mathbf{S}_t^\dagger$ fall in the shaded region on the complex plane, then the inequality (3) is satisfied.

Canceling $\mathbf{W}_{t+1}\mathbf{A}_h$ gives us

$$(\mathbf{S}_{t+1} - \mathbf{U}_{t+1})\mathbf{S}_h^\dagger = \mathbf{S}_t\mathbf{S}_h^\dagger + \gamma(\mathbf{S}_t\mathbf{S}_h^\top)^{-\top}. \quad (5)$$

Since we assume $\mathbf{S}_t\mathbf{S}_h^\dagger$, so is $(\mathbf{S}_t\mathbf{S}_h^\top)^{-\top}$, and noticing that they share the same eigen basis, all eigenvalues of (5) satisfies (4), therefore we have

$$\log |\det(\mathbf{S}_{t+1} - \mathbf{U}_{t+1})\mathbf{S}_h^\dagger| \leq \text{Tr}(\mathbf{S}_{t+1} - \mathbf{U}_{t+1})\mathbf{S}_h^\dagger - k. \quad (6)$$

Notice that

$$\mathbf{S}_{t+1} - \mathbf{U}_{t+1} = \mathbf{W}_{t+1}\mathbf{X} - \mathbf{U}_t,$$

while

$$\mathbf{S}_{t+1} = \text{Prox}_{\gamma g}(\mathbf{W}_{t+1}\mathbf{X} - \mathbf{U}_t),$$

so we equivalently have

$$\mathbf{S}_{t+1} = \text{Prox}_{\gamma g}(\mathbf{S}_{t+1} - \mathbf{U}_{t+1}).$$

Since we have (6), by using the property of proximal operator, we have

$$\log |\det \text{Prox}_{\gamma g}(\mathbf{S}_{t+1} - \mathbf{U}_{t+1})\mathbf{S}_h^\dagger| \leq \text{Tr} \text{Prox}_{\gamma g}(\mathbf{S}_{t+1} - \mathbf{U}_{t+1})\mathbf{S}_h^\dagger - k.$$

This completes the proof. $\square$

C PCA EQUIVALENCE

Consider the main formulation with $g(\cdot) = (1/2)\|\cdot\|_F^2$:

$$\underset{\mathbf{W}}{\text{minimize}} \quad -\frac{1}{2} \log \det \mathbf{W}\mathbf{W}^\top + \frac{1}{2} \|\mathbf{W}\mathbf{X}\|_F^2. \quad (7)$$

Taking the gradient with respect to $\mathbf{W}$ and setting it equal to zero at optimum $\mathbf{W}_\star$, we have

$$\mathbf{W}_\star^{\dagger\top} = \mathbf{W}_\star\mathbf{X}\mathbf{X}^\top.$$

We assume that $\mathbf{W}_\star$ is a wide matrix with linearly independent rows, so $\mathbf{W}_\star\mathbf{W}_\star^\dagger = \mathbf{I}$, therefore multiplying both sides by $\mathbf{W}_\star^\top$ on the right gets

$$\mathbf{I} = \mathbf{W}_\star\mathbf{X}\mathbf{X}^\top\mathbf{W}_\star^\top$$

This shows two things:

1. The symmetric matrix $\mathbf{X}\mathbf{X}^\top$ can be diagonalized by $\mathbf{W}_\star$. This can happen if rows of $\mathbf{W}_\star$ are eigenvectors of $\mathbf{X}\mathbf{X}^\top$ and their Euclidean norms equal to one over the square root of their corresponding eigenvalue.
2. At optimum, the second term in (7) exactly equals to $(1/2)\|\mathbf{W}_\star\mathbf{X}\|_F^2 = k/2$, meaning that we only need to worry about the first term when picking the eigenvectors of $\mathbf{X}\mathbf{X}^\top$ to construct $\mathbf{W}_\star$.

As a result, it is easy to conclude that an optimal solution is $\mathbf{W}_\star = \mathbf{\Sigma}_k^{-1}\mathbf{U}_k^\top$, where $\mathbf{\Sigma}$ is a diagonal matrix with the k largest singular values of $\mathbf{X}$ on the diagonal, and columns of $\mathbf{U}_k$ are their corresponding left singular vectors. Note that $\mathbf{Q}\mathbf{W}_\star$ would also be optimal for any orthogonal matrix $\mathbf{Q}$, showing that this model is not identifiable.

This indeed provides yet another interpretation of the famous PCA problem.

REFERENCES

- Dimitri P. Bertsekas. *Nonlinear Programming*. Athena Scientific, 2 edition, 1999.
- José M Bioucas-Dias. A variable splitting augmented lagrangian approach to linear spectral unmixing. In *2009 First workshop on hyperspectral image and signal processing: Evolution in remote sensing*, pp. 1–4. IEEE, 2009.

- Tsung-Han Chan, Chong-Yung Chi, Yu-Min Huang, and Wing-Kin Ma. A convex analysis-based minimum-volume enclosing simplex algorithm for hyperspectral unmixing. *IEEE Transactions on Signal Processing*, 57(11):4418–4432, 2009.
- Jingzhou Hu and Kejun Huang. Global identifiability of ℓ_1 -based dictionary learning via matrix volume optimization. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 36, 2023a.
- Jingzhou Hu and Kejun Huang. Identifiable bounded component analysis via minimum volume enclosing parallelotope. In *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2023b.
- Jingzhou Hu and Kejun Huang. Frank-wolfe algorithm for simplicial and nonnegative component analysis. In *2024 IEEE 13rd Sensor Array and Multichannel Signal Processing Workshop (SAM)*, pp. 1–5. IEEE, 2024.
- Kejun Huang and Xiao Fu. Detecting overlapping and correlated communities without pure nodes: Identifiability and algorithm. In *International Conference on Machine Learning*, pp. 2859–2868, 2019.
- Kejun Huang, Xiao Fu, and Nikolaos D Sidiropoulos. Anchor-free correlated topic modeling: Identifiability and algorithm. *Advances in Neural Information Processing Systems*, 29, 2016.