

An Iterative Algorithm for Regularized Non-negative Matrix Factorizations

Steven E. Pav *

July 14, 2026

Abstract

We generalize the non-negative matrix factorization algorithm of Lee and Seung [2] to accept a weighted norm, and to support ridge and Lasso regularization. We recast the Lee and Seung multiplicative update as an additive update which does not get stuck on zero values.

1 Introduction

The term *non-negative matrix factorization* refers to decomposing a given matrix Y with non-negative elements approximately as the product $Y \approx LR$, for conformable non-negative matrices L and R of smaller rank than Y . This factorization represents a kind of dimensionality reduction. The quality of the approximation is measured by some “divergence” between Y and LR , or some norm of their difference. Perhaps the most common is the Frobenius (or elementwise ℓ_2) distance between them. Under this objective, non-negative matrix factorization can be expressed as the optimization problem:

$$\min_{L \geq 0, R \geq 0} \text{tr} \left((Y - LR)^\top (Y - LR) \right), \quad (1)$$

where $\text{tr}(M)$ is the trace of the matrix M , and $A \geq 0$ means the elements of A are non-negative. [8]

Lee and Seung described an iterative approach to this problem, as well as for a related optimization with a different objective. [2] Their algorithm starts with some unspecified initial iterates $L^{(0)}$ and $R^{(0)}$, then iteratively improves $L^{(k)}$ and $R^{(k)}$ in turn. That is, given $L^{(k)}$ and $R^{(k)}$, an update is made to arrive at $L^{(k+1)}$; then using $L^{(k+1)}$ and $R^{(k)}$, one computes $R^{(k+1)}$, and the process repeats. The iterative update for $L^{(k+1)}$ is

$$L^{(k+1)} \leftarrow L^{(k)} \odot Y R^{(k)\top} \oslash \left(L^{(k)} R^{(k)} R^{(k)\top} \right). \quad (2)$$

We use $\odot$ to mean the Hadamard, or ‘elementwise’, multiplication, and $\oslash$ to mean Hadamard division. The iterative update for $R^{(k+1)}$ is defined *mutatis mutandis*. As should be apparent from inspection, this iterative update maintains non-negativity of estimates. That is, if Y , $L^{(k)}$, and $R^{(k)}$ have only positive

* shabbychef@gmail.com Application of the multiplicative methods described herein for chemometrics applications may be covered by U.S. Patent. [6]

elements, then $\mathbf{L}^{(k+1)}$ is finite and has only positive elements. It is not clear what should be done if elements of the denominator term $\mathbf{L}^{(k)} \mathbf{R}^{(k)} \mathbf{R}^{(k)\top}$ are zero, and whether and how this can be avoided.

Gonzalez and Zhang describe a modification of the update to accelerate convergence. [1] Their accelerated algorithm updates $\mathbf{L}^{(k)}$ row by row, and $\mathbf{R}^{(k)}$ column by column, using the step length modification of Merritt and Zhang [4].

One attractive feature of the Lee and Seung iterative update is its sheer simplicity. In a high-level computer language with a rich set of matrix operations, such as Julia or octave, the iterative step can be expressed in a single line of code; the entire algorithm could be expressed in perhaps 20 lines of code. Conceivably this algorithm could even be implemented in challenged computing environments, like within a SQL query or a spreadsheet macro language. Simplicity makes this algorithm an obvious choice for experimentation without requiring an all-purpose constrained quadratic optimization solver, or other complicated software.

2 Regularized Non-Negative Matrix Factorization

We generalize the optimization problem of 1 to include a weighted norm, regularization terms, and orthogonality penalties. For a given $\mathbf{Y}$, and non-negative weighting matrices $\mathbf{W}_{0,R}$, $\mathbf{W}_{0,C}$, $\mathbf{W}_{1,L}$, $\mathbf{W}_{1,R}$, $\mathbf{W}_{2,R,L,j}$, $\mathbf{W}_{2,C,L,j}$, $\mathbf{W}_{2,R,R,j}$, and $\mathbf{W}_{2,C,R,j}$, find matrices $\mathbf{L}$ and $\mathbf{R}$ to minimize

$$\begin{aligned} \min_{\mathbf{L} \geq 0, \mathbf{R} \geq 0} \quad & \frac{1}{2} \text{tr} \left((\mathbf{Y} - \mathbf{LR})^\top \mathbf{W}_{0,R} (\mathbf{Y} - \mathbf{LR}) \mathbf{W}_{0,C} \right) \\ & + \text{tr} \left(\mathbf{W}_{1,L}^\top \mathbf{L} \right) + \text{tr} \left(\mathbf{W}_{1,R}^\top \mathbf{R} \right) \\ & + \frac{1}{2} \sum_{1 \leq j \leq J} \text{tr} \left(\mathbf{L}^\top \mathbf{W}_{2,R,L,j} \mathbf{L} \mathbf{W}_{2,C,L,j} \right) \\ & + \frac{1}{2} \sum_{1 \leq j \leq J} \text{tr} \left(\mathbf{R}^\top \mathbf{W}_{2,R,R,j} \mathbf{R} \mathbf{W}_{2,C,R,j} \right). \end{aligned} \quad (3)$$

Essentially we are trying to model $\mathbf{Y}$ as $\mathbf{LR}$. Our loss function is quadratic in the error $\mathbf{Y} - \mathbf{LR}$, with weights $\mathbf{W}_{0,R}$ and $\mathbf{W}_{0,C}$, and ℓ_1 and ℓ_2 regularization via the terms with $\mathbf{W}_1$ and $\mathbf{W}_{2,R,j}$ and $\mathbf{W}_{2,C,j}$. This optimization problem includes that of 1 as a special case.

We note that a non-orthogonality penalty term on the $\mathbf{X}$ can be expressed as

$$\frac{c}{2} \text{tr} \left(\mathbf{X}^\top \mathbf{IX} (\mathbf{11}^\top - \mathbf{I}) \right).$$

Thus a non-orthogonality penalty term can be expressed with the $\mathbf{W}_{2,R,j}$ and $\mathbf{W}_{2,C,j}$. Indeed the motivation for including a sum of the ℓ_2 terms is to allow for both straightforward regularization, where $\mathbf{W}_{2,R,j} = \mathbf{I}$ and $\mathbf{W}_{2,C,j} = c\mathbf{I}$, and orthogonality terms. For simplicity of notation, in the discussion that follows we will often treat the ℓ_2 terms as if there were only one of them, *i.e.*, $J = 1$, and omit the j subscripts. Restoring the more general case consists only of expanding to a full sum of J terms.

This problem formulation is appropriate for power users who can meaningfully select the various weighting matrices. And while it is easy to conceive of a use for general forms for $W_{0,R}$ and $W_{0,C}$ (the rows of Y have different importances, or the columns of Y are observed with different accuracies, for example), most users would probably prefer simpler formulations for the $W_{1,\cdot}$ and $W_{2,\cdot,\cdot,j}$ matrices. To achieve something like an elasticnet factorization with non-orthogonality penalties, one should set $J = 2$ and take

$$\begin{aligned} W_{1,L} &= \lambda_{1,L} \mathbf{1}, & W_{1,R} &= \lambda_{1,R} \mathbf{1}, \\ W_{2,R,L,1} &= \lambda_{2,L} \mathbf{I}, & W_{2,C,L,1} &= \mathbf{I}, \\ W_{2,R,R,1} &= \lambda_{2,R} \mathbf{I}, & W_{2,C,R,1} &= \mathbf{I}, \\ W_{2,R,L,2} &= \mathbf{I}, & W_{2,C,L,2} &= \gamma_{2,L} (\mathbf{1}\mathbf{1}^\top - \mathbf{I}), \\ W_{2,R,R,2} &= \mathbf{I}, & W_{2,C,R,2} &= \gamma_{2,R} (\mathbf{1}\mathbf{1}^\top - \mathbf{I}), \end{aligned} \tag{4}$$

where here $\mathbf{1}$ stands in for an all-ones matrix of the appropriate size. All the λ and γ parameters must be non-negative, and the identity matrices are all appropriately sized. This is the reduced form of the problem that depends mostly on a few scalars. A simple form for the weighted error is to take $W_{0,R}$ and $W_{0,C}$ to be diagonal matrices,

Depending on the regularization terms, the solution to the problem of 3 is unlikely to be unique. Certain the solution to 1 is not unique, since if Q is a permutation matrix of the appropriate size, then LQ and $Q^{-1}R$ are another set of non-negative matrices with the same objective value. Perhaps the regularization terms can help avoid these kinds of ambiguity in the solution.

2.0.1 Split Form

To approximately solve this problem our algorithm starts with some initial estimates of L and R and iteratively alternates between updating the estimate of L with the estimate of R fixed, and updating the estimate of R with the estimate of L fixed. To discuss these half-step updates, we can describe both of them in a single unified formulation. So we consider the optimization problem

$$\min_{X \geq 0} \phi(X), \tag{5}$$

where we define the objective function

$$\begin{aligned} \phi(X) &= \frac{1}{2} \text{tr} \left((Y - LXR)^\top W_{0,R} (Y - LXR) W_{0,C} \right) \\ &\quad + \text{tr} (W_1^\top X) + \frac{1}{2} \sum_{1 \leq j \leq J} \text{tr} (X^\top W_{2,R,j} X W_{2,C,j}). \end{aligned} \tag{6}$$

We must assume that L has full column rank and R has full row rank, $W_{0,R}$, $W_{0,C}$, $W_{2,R,j}$ and $W_{2,C,j}$ are square, symmetric and positive semidefinite. The $W_{0,R}$ has the same number of rows as Y , and $W_{0,C}$ has the same number of columns as Y . W_1 is the same size as X ; we will make further restrictions on W_1 below. The matrices $W_{2,R,j}$ have the same number of rows as X , and $W_{2,C,j}$ have the same number of columns as X . All the $W_{i,\cdot}$ are assumed to have non-negative elements, that is $W_{i,\cdot} \geq 0$.

Our algorithm now consists of filling in an identity for $\mathbf{L}$, and estimate of $\mathbf{L}$ for $\mathbf{X}$, and all the $\mathbf{L}$ -appropriate weighting matrices into Equation 6 and taking a step to solve Problem 3, getting an update of $\mathbf{L}$. Then we perform a similar operation to get a new estimate of $\mathbf{R}$.

2.1 Factorization in Vector Form

The analysis is simplified if we express the problem as one of finding a vector unknown, so we will ‘vectorize’ the problem. For matrix $\mathbf{M}$ let $\text{vec}(\mathbf{M})$ be the vector which consists of the columns of $\mathbf{M}$ stacked on top of each other; let $\mathbf{A} \otimes \mathbf{B}$ denote the Kronecker product of $\mathbf{A}$ and $\mathbf{B}$. [9] See Section A in the appendix for some helpful identities involving the Kronecker product, matrix trace, and vectorization.

We write $\phi(\mathbf{X})$ as a quadratic function of the vector $\text{vec}(\mathbf{X})$

$$\begin{aligned}\phi(\mathbf{X}) &= \frac{1}{2} \|\text{vec}(\mathbf{Y}) - (\mathbf{R}^\top \otimes \mathbf{L}) \text{vec}(\mathbf{X})\|_{\mathbf{W}_{0,C} \otimes \mathbf{W}_{0,R}}^2 \\ &\quad + \frac{1}{2} \sum_j \|\text{vec}(\mathbf{X})\|_{\mathbf{W}_{2,C,j} \otimes \mathbf{W}_{2,R,j}}^2 + \text{vec}(\mathbf{W}_1)^\top \text{vec}(\mathbf{X}) \\ &= \frac{1}{2} \text{vec}(\mathbf{X})^\top \mathbf{G} \text{vec}(\mathbf{X}) + \mathbf{d}^\top \text{vec}(\mathbf{X}) + c,\end{aligned}\tag{7}$$

where

$$\begin{aligned}\mathbf{G} &= (\mathbf{R}^\top \otimes \mathbf{L})^\top (\mathbf{W}_{0,C} \otimes \mathbf{W}_{0,R}) (\mathbf{R}^\top \otimes \mathbf{L}) + \sum_j \mathbf{W}_{2,C,j} \otimes \mathbf{W}_{2,R,j}, \\ &= (\mathbf{R} \mathbf{W}_{0,C} \mathbf{R}^\top) \otimes (\mathbf{L}^\top \mathbf{W}_{0,R} \mathbf{L}) + \sum_j \mathbf{W}_{2,C,j} \otimes \mathbf{W}_{2,R,j}. \\ \mathbf{d} &= \text{vec}(\mathbf{W}_1) - (\mathbf{R}^\top \otimes \mathbf{L})^\top (\mathbf{W}_{0,C} \otimes \mathbf{W}_{0,R}) \text{vec}(\mathbf{Y}), \\ &= \text{vec}(\mathbf{W}_1) - \text{vec}(\mathbf{L}^\top \mathbf{W}_{0,R} \mathbf{Y} \mathbf{W}_{0,C} \mathbf{R}^\top).\end{aligned}\tag{8}$$

We write $\|\mathbf{x}\|_{\mathbf{A}}^2$ to mean $\mathbf{x}^\top \mathbf{A} \mathbf{x}$. Note that $\mathbf{G} \geq 0$, but $\mathbf{d}$ will likely have negative elements. In fact, to use Lee and Seung’s iterative algorithm, the vector $\mathbf{d}$ must have non-positive elements. This is equivalent to the restriction $\mathbf{L}^\top \mathbf{W}_{0,R} \mathbf{Y} \mathbf{W}_{0,C} \mathbf{R} - \mathbf{W}_1 \geq 0$, taken elementwise. This imposes a restriction on $\mathbf{W}_1$ which is hard to verify for the problem of non-negative matrix factorization. Note that by definition $\mathbf{G}$ can be positive definite, depending on the $\mathbf{R}$, $\mathbf{L}$ and the $\mathbf{W}_{2,\cdot}$, but is only positive *semi* definite in the general case.

We can thus consider this problem as optimization over a vector unknown. In an abuse of notation, we reuse the objective function $\phi(\cdot)$, writing

$$\phi(\mathbf{x}) = \frac{1}{2} \mathbf{x}^\top \mathbf{G} \mathbf{x} + \mathbf{d}^\top \mathbf{x},\tag{9}$$

for symmetric, nonnegative positive semidefinite $\mathbf{G}$. So now we consider the positivity-constrained quadratic optimization problem

$$\min_{\mathbf{x} \geq 0} \phi(\mathbf{x}).\tag{10}$$

Note that the unconstrained solution to this problem is $-\mathbf{G}^{-1} \mathbf{d}$. When $\mathbf{G}$ is positive definite, because the constraint space is convex, this solution will be

unique. [5] In general, however, we will not have unique solutions. When the W_1 term is all zero, this a weighted least squares problem, possibly with Tikhonov regularization.

The gradient of the objective $\phi(\mathbf{x})$ is $\nabla \phi(\mathbf{x}) = \mathbf{G}\mathbf{x} + \mathbf{d}$. When considering the matrix form of the problem, this has the value

$$\text{vec} \left(\left[\mathbf{L}^\top \mathbf{W}_{0,R} \mathbf{L} \mathbf{X} \mathbf{R} \mathbf{W}_{0,C} \mathbf{R}^\top + \sum_j \mathbf{W}_{2,R,j} \mathbf{X} \mathbf{W}_{2,C,j} + \mathbf{W}_1 - \mathbf{L}^\top \mathbf{W}_{0,R} \mathbf{Y} \mathbf{W}_{0,C} \mathbf{R}^\top \right] \right).$$

We consider an iterative solution to this problem. We first select some $\mathbf{x}^{(0)}$. Then given current iterate $\mathbf{x}^{(k)}$, we seek to find the next iterate $\mathbf{x}^{(k+1)}$. Define

$$\phi^{(k)}(\mathbf{x}) = \frac{1}{2} \mathbf{x}^\top \mathbf{G}^{(k)} \mathbf{x} + \mathbf{d}^{(k)\top} \mathbf{x} + c_k,$$

where we will set $\mathbf{G}^{(k)}$, $\mathbf{d}^{(k)}$ and c_k such that $\phi^{(k)}(\mathbf{x})$ is tangent to $\phi(\mathbf{x})$ at $\mathbf{x}^{(k)}$, and such that $\phi^{(k)}(\mathbf{x}) \geq \phi(\mathbf{x})$ for all $\mathbf{x}$. We will take $\mathbf{x}^{(k+1)}$ as some point which improves $\phi^{(k)}(\mathbf{x})$: its minimizer if it is non-negative, or another point if not. If $\mathbf{x}^{(k+1)}$ improves $\phi^{(k)}(\mathbf{x})$, then it also improves $\phi(\mathbf{x})$, as

$$\phi(\mathbf{x}^{(k+1)}) \leq \phi^{(k)}(\mathbf{x}^{(k+1)}) \leq \phi^{(k)}(\mathbf{x}^{(k)}) = \phi(\mathbf{x}^{(k)}).$$

The tangency condition fixes the value of c_k , and also gives the identity

$$\mathbf{d}^{(k)} = \mathbf{d} + \mathbf{G}\mathbf{x}^{(k)} - \mathbf{G}^{(k)}\mathbf{x}^{(k)}.$$

To ensure that $\phi^{(k)}(\mathbf{x}) \geq \phi(\mathbf{x})$ it is sufficient to select $\mathbf{G}^{(k)}$ such that $\mathbf{G}^{(k)} - \mathbf{G}$ is positive semidefinite. The following lemma gives us a way of selecting diagonal matrices $\mathbf{G}^{(k)}$ that satisfy this condition.

First a bit of notation, for vector $\mathbf{v}$, $\text{Diag}(\mathbf{v})$ is the diagonal matrix with $\mathbf{v}$ for diagonal. For matrix $\mathbf{A}$, $\text{diag}(\mathbf{A})$ is the vector of the diagonal of $\mathbf{A}$. Thus we have $\text{diag}(\text{Diag}(\mathbf{v})) = \mathbf{v}$.

Lemma 2.1. *Let $\mathbf{G}$ be a symmetric matrix with non-negative elements, and full rank, and let $\mathbf{b}$ be a vector with strictly positive elements. Then*

$$\text{Diag}(\mathbf{G}\mathbf{b}) \text{Diag}(\mathbf{b})^{-1} \succeq \mathbf{G}.$$

(Here $\mathbf{A} \succeq \mathbf{B}$ means that $\mathbf{A} - \mathbf{B}$ is positive semidefinite.)

Proof. First note that if $\mathbf{A}$ is symmetric, and has non-negative elements, then $\text{Diag}(\mathbf{A}\mathbf{1}) \succeq \mathbf{A}$ because the matrix $\text{Diag}(\mathbf{A}\mathbf{1}) - \mathbf{A}$ is symmetric, with non-negative diagonal, and is diagonally dominant, thus it is positive semidefinite.

Apply that fact with $\mathbf{A} = \text{Diag}(\mathbf{b}) \mathbf{G} \text{Diag}(\mathbf{b})$. Then

$$\begin{aligned} \text{Diag}(\text{Diag}(\mathbf{b}) \mathbf{G} \text{Diag}(\mathbf{b}) \mathbf{1}) &\succeq \text{Diag}(\mathbf{b}) \mathbf{G} \text{Diag}(\mathbf{b}), \\ \text{Diag}(\mathbf{b}) \text{Diag}(\mathbf{G}\mathbf{b}) &\succeq \text{Diag}(\mathbf{b}) \mathbf{G} \text{Diag}(\mathbf{b}), \\ \text{Diag}(\mathbf{G}\mathbf{b}) \text{Diag}(\mathbf{b})^{-1} &\succeq \mathbf{G}, \end{aligned}$$

as needed.

Note that the proof relies on strict positivity of the elements of $\mathbf{b}$. By this we can claim a bijection between $\mathbb{R}^n$ and $\text{Diag}(\mathbf{b}) \mathbb{R}^n$, which allows us to conclude the last line above. $\square$

Lemma 2.2. *Let $\mathbf{G}$ be a symmetric matrix with non-negative elements, and full rank. Let $\mathbf{b}$ be some vector with non-negative elements such that $\mathbf{G}\mathbf{b}$ has strictly positive elements. Letting*

$$\mathbf{x}^{(k+1)} \leftarrow \mathbf{x}^{(k)} - \text{Diag}(\mathbf{b}) \text{Diag}(\mathbf{G}\mathbf{b})^{-1} (\mathbf{d} + \mathbf{G}\mathbf{x}^{(k)}),$$

then $\phi(\mathbf{x}^{(k+1)}) \leq \phi(\mathbf{x}^{(k)})$.

Proof. First, consider the case where $\mathbf{b}$ has strictly positive elements. Letting $\mathbf{G}^{(k)} = \text{Diag}(\mathbf{G}\mathbf{b}) \text{Diag}(\mathbf{b})^{-1}$, the minimum of $\phi^{(k)}(\mathbf{x})$ occurs at $-\mathbf{G}^{(k)-1}\mathbf{d}^{(k)}$, which has value

$$\begin{aligned} -\mathbf{G}^{(k)-1}\mathbf{d}^{(k)} &= -\mathbf{G}^{(k)-1} (\mathbf{d} + \mathbf{G}\mathbf{x}^{(k)} - \mathbf{G}^{(k)}\mathbf{x}^{(k)}), \\ &= \mathbf{x}^{(k)} - \mathbf{G}^{(k)-1} (\mathbf{d} + \mathbf{G}\mathbf{x}^{(k)}), \\ &= \mathbf{x}^{(k)} - \text{Diag}(\mathbf{b}) \text{Diag}(\mathbf{G}\mathbf{b})^{-1} (\mathbf{d} + \mathbf{G}\mathbf{x}^{(k)}) = \mathbf{x}^{(k+1)}. \end{aligned} \quad (11)$$

By Lemma 2.1, $\phi^{(k)}(\mathbf{x})$ dominates $\phi(\mathbf{x})$, and so $\phi(\mathbf{x}^{(k+1)}) \leq \phi(\mathbf{x}^{(k)})$.

To prove the theorem for the general case where $\mathbf{b}$ is simply non-negative, consider a sequence of strictly positive vectors which converge to $\mathbf{b}$, and apply the argument above, then appeal to continuity of $\phi(\cdot)$. $\square$

Theorem 2.3 (Lee and Seung). *Let $\mathbf{G}$ be a symmetric matrix with non-negative elements, and full rank. Assume $\mathbf{x}^{(k)}$ has non-negative elements, and $\mathbf{G}\mathbf{x}^{(k)}$ has strictly positive elements. If*

$$\mathbf{x}^{(k+1)} \leftarrow -\mathbf{x}^{(k)} \odot \mathbf{d} \oslash (\mathbf{G}\mathbf{x}^{(k)}), \quad (12)$$

then $\phi(\mathbf{x}^{(k+1)}) \leq \phi(\mathbf{x}^{(k)})$. Moreover, the update preserves non-negativity of iterates as long as $\mathbf{d}$ has non-positive elements.

Proof. Let $\mathbf{b} = \mathbf{x}^{(k)}$ in Lemma 2.2. Then note that

$$\begin{aligned} \mathbf{x}^{(k+1)} &= \mathbf{x}^{(k)} - \text{Diag}(\mathbf{x}^{(k)}) \text{Diag}(\mathbf{G}\mathbf{x}^{(k)})^{-1} (\mathbf{d} + \mathbf{G}\mathbf{x}^{(k)}), \\ &= \mathbf{x}^{(k)} - \text{Diag}(\mathbf{x}^{(k)}) \text{Diag}(\mathbf{G}\mathbf{x}^{(k)})^{-1} \mathbf{d} - \mathbf{x}^{(k)}, \\ &= -\text{Diag}(\mathbf{x}^{(k)}) \text{Diag}(\mathbf{G}\mathbf{x}^{(k)})^{-1} \mathbf{d}, \\ &= -\mathbf{x}^{(k)} \odot \mathbf{d} \oslash (\mathbf{G}\mathbf{x}^{(k)}). \end{aligned} \quad (13)$$

$\square$

Returning to the original problem of minimizing the objective of Equation 6, the iterative update is

$$\begin{aligned} \mathbf{X}^{(k+1)} &\leftarrow -\mathbf{X}^{(k)} \odot (\mathbf{W}_1 - \mathbf{L}^\top \mathbf{W}_{0,R} \mathbf{Y} \mathbf{W}_{0,C} \mathbf{R}^\top) \oslash \\ &\quad \left(\mathbf{L}^\top \mathbf{W}_{0,R} \mathbf{L} \mathbf{X}^{(k)} \mathbf{R} \mathbf{W}_{0,C} \mathbf{R}^\top + \sum_j \mathbf{W}_{2,R,j} \mathbf{X}^{(k)} \mathbf{W}_{2,C,j} \right). \end{aligned} \quad (14)$$

2.1.1 As a Matrix Factorization Algorithm

Now consider using this iterative update in the solution of Problem 3. One starts with initial guesses $\mathbf{L}^{(0)}$ and $\mathbf{R}^{(0)}$, which are strictly positive, then computes in turn

$$\mathbf{L}^{(1)} \leftarrow -\mathbf{L}^{(0)} \odot \left(\mathbf{W}_{1,L} - \mathbf{W}_{0,R} \mathbf{Y} \mathbf{W}_{0,C} \mathbf{R}^{(0)\top} \right) \oslash \quad (15)$$

$$\left(\mathbf{W}_{0,R} \mathbf{L}^{(0)} \mathbf{R}^{(0)} \mathbf{W}_{0,C} \mathbf{R}^{(0)\top} + \sum_j \mathbf{W}_{2,R,L,j} \mathbf{L}^{(0)} \mathbf{W}_{2,C,L,j} \right),$$

$$\mathbf{R}^{(1)} \leftarrow -\mathbf{R}^{(0)} \odot \left(\mathbf{W}_{1,R} - \mathbf{L}^{(1)\top} \mathbf{W}_{0,R} \mathbf{Y} \mathbf{W}_{0,C} \right) \oslash \quad (16)$$

$$\left(\mathbf{L}^{(1)\top} \mathbf{W}_{0,R} \mathbf{L}^{(1)} \mathbf{R}^{(0)} \mathbf{W}_{0,C} + \sum_j \mathbf{W}_{2,R,R,j} \mathbf{R}^{(0)} \mathbf{W}_{2,C,R,j} \right).$$

Then one computes estimates of $\mathbf{L}^{(2)}, \mathbf{R}^{(2)}, \mathbf{L}^{(3)}, \mathbf{R}^{(3)}, \dots$

The restriction on the $\mathbf{W}_1$ from the symmetric form of the problem translates into the requirements that

$$\mathbf{W}_{1,L} - \mathbf{W}_{0,R} \mathbf{Y} \mathbf{W}_{0,C} \mathbf{R}^{(k)\top} \leq 0, \quad \text{and} \quad (17)$$

$$\mathbf{W}_{1,R} - \mathbf{L}^{(k)\top} \mathbf{W}_{0,R} \mathbf{Y} \mathbf{W}_{0,C} \leq 0, \quad (18)$$

elementwise. This restriction suffices to keep iterates strictly positive if the initial iterates are also strictly positive. Note, however, it is not clear how this can be guaranteed *a priori*, as the iterates $\mathbf{R}^{(k)}$ and $\mathbf{L}^{(k)}$ may become small, and the $\mathbf{W}_{1,\cdot}$ may be large. In fact, the *entire point* of ℓ_1 regularization is to encourage elements of $\mathbf{R}^{(k)}$ and $\mathbf{L}^{(k)}$ to take value zero. As a practical matter, then, code which allows arbitrary $\mathbf{W}_{1,L}, \mathbf{W}_{1,R}$ should take the liberty of overriding the input values and always guaranteeing that the numerators are bounded elementwise by

$$\left(\mathbf{W}_{1,L} - \mathbf{W}_{0,R} \mathbf{Y} \mathbf{W}_{0,C} \mathbf{R}^{(k)\top} \right) \leq -\epsilon \geq \left(\mathbf{W}_{1,R} - \mathbf{L}^{(k)\top} \mathbf{W}_{0,R} \mathbf{Y} \mathbf{W}_{0,C} \right) \quad (19)$$

for some ϵ . This is a compromise for simplicity, and thus this form of the algorithm does not always solve for the input problem. In the next section we describe a more principled approach that respects the user input.

Again the restriction on the $\mathbf{W}_1$ from the symmetric problem cannot be guaranteed *a priori*; instead an implementation must keep the numerators bounded away from zero. We present the method in Algorithm 1. We call it a “multiplicative update” since the update steps are elementwise multiply and divide operations on the initial estimates of $\mathbf{L}$ and $\mathbf{R}$. For this reason, it has the same limitations as the Lee and Seung algorithm, namely that once an element of $\mathbf{L}$ or $\mathbf{R}$ takes value zero, it can never take a non-zero value. This algorithm assumes the simplified form of Equations 4; the code for the more general form of the problem is an obvious modification of the update steps.

3 Additive Steps and Convergence

Theorem 2.3 only tells us that the sequence $\phi(\mathbf{x}^{(k)})$ is non-decreasing. It does *not* guarantee convergence to a global, or even a local, minimum. In fact, if we

Algorithm 1 Multiplicative Update Regularized Non-Negative Matrix Factorization.

function MURNMF($Y, d, W_{0,R}, W_{0,C}, \lambda_{1,L}, \lambda_{1,R}, \lambda_{2,L}, \lambda_{2,R}, \gamma_{2,L}, \gamma_{2,R}, \epsilon > 0$)

 Initialize random matrices $L^{(0)} > 0$ and $R^{(0)} > 0$ with d columns and rows, respectively.

 Let $k \leftarrow 0$.

while not converged **do**

 Compute numerator $H \leftarrow (\lambda_{1,L} \mathbf{1} - W_{0,R} Y W_{0,C} R^{(k)\top})$.

 Clip the numerator $H \leftarrow H \wedge -\epsilon \mathbf{1}$.

 Let $F \leftarrow (W_{0,R} L^{(k)} R^{(k)} W_{0,C} R^{(k)\top} + \lambda_{2,L} L^{(k)} + \gamma_{2,L} L^{(k)} (\mathbf{1} \mathbf{1}^\top - I))$.

 Let $L^{(k+1)} \leftarrow -L^{(k)} \odot H \oslash F$.

 Compute numerator $J \leftarrow (\lambda_{1,R} \mathbf{1} - L^{(k+1)\top} W_{0,R} Y W_{0,C})$.

 Clip the numerator $J \leftarrow J \wedge -\epsilon \mathbf{1}$.

 Let $F \leftarrow (L^{(k+1)\top} W_{0,R} L^{(k+1)} R^{(k)} W_{0,C} + \lambda_{2,R} R^{(k)} + \gamma_{2,R} R^{(k)} (\mathbf{1} \mathbf{1}^\top - I))$.

 Let $R^{(k+1)} \leftarrow -R^{(k)} \odot J \oslash F$.

 Increment $k \leftarrow k + 1$.

 Check convergence.

end while

return $x^{(k)}$

end function

only restrict d to be non-positive, it is easy to construct cases where $x^{(k)}$ will not converge to the constrained optimum. For if an element of d is zero, then that same element of $x^{(k)}$ will be zero for all $k > 0$. However, the optimum may occur for x with a non-zero value for that element. (Certainly the *unconstrained* solution $-G^{-1}d$ may have a non-zero value for that element.)

Here we analyze the Lee and Seung update as a more traditional additive step, rather than a multiplicative step. Expressing the iteration as an additive step, we have

$$x^{(k+1)} = x^{(k)} + h^{(k+1)}, \quad (20)$$

where

$$h^{(k+1)} = -x^{(k)} \odot (d \oslash (Gx^{(k)} + \mathbf{1})), \quad (21)$$

$$\begin{aligned} &= - (Gx^{(k)} + d) \odot x^{(k)} \oslash (Gx^{(k)}), \\ &= -\nabla \phi(x^{(k)}) \odot x^{(k)} \oslash (Gx^{(k)}). \end{aligned} \quad (22)$$

This may not be the optimal length step in the direction of $h^{(k+1)}$. If we instead take

$$x^{(k+1)} = x^{(k)} + ch^{(k+1)}, \quad (23)$$

then $\phi(x^{(k+1)})$ is a quadratic function of c with optimum at

$$c^* = \frac{-(Gx^{(k)} + d)^\top h^{(k+1)}}{h^{(k+1)\top} G h^{(k+1)}}. \quad (24)$$

However, if $c^* > 1$ elements of $\mathbf{x}^{(k)} + c^* \mathbf{h}^{(k+1)}$ could be negative. One would instead have to select a c that results in a strictly non-negative $\mathbf{x}^{(k+1)}$. However, if $c^* < 1$, then the original algorithm would have overshoot the optimum.

Lemma 3.1. *Let $\mathbf{G}$ be a symmetric matrix with non-negative elements, and full rank. Let c^* be as defined in Equation 24. Let c' be the largest number in $[0, c^*]$ such that all elements of $\mathbf{x}^{(k+1)} \leftarrow \mathbf{x}^{(k)} + c' \mathbf{h}^{(k+1)}$ are non-negative, where $\mathbf{h}^{(k+1)}$ is as in Equation 22. Then $\phi(\mathbf{x}^{(k+1)}) \leq \phi(\mathbf{x}^{(k)})$.*

Proof. Letting $f(c) = \phi(\mathbf{x}^{(k)} + c \mathbf{h}^{(k+1)})$, $f(c)$ is quadratic in c with positive second derivative. Then since c' is between 0 and c^* , we have $f(c') \leq f(0)$, which was to be proven. By construction the elements of $\mathbf{x}^{(k+1)}$ are non-negative. $\square$

We note that if $c^* \leq 1$ then $c' = c^*$, since non-negativity is sustained for smaller step sizes.

This iterative update was originally described by Merritt and Zhang for solution of “Totally Nonnegative Least Squares” problems. [4] Their algorithm has guaranteed convergence for minimization of $\|\mathbf{A}\mathbf{x} - \mathbf{b}\|_2^2$, without the weighting or regularization terms, under certain conditions. They assume strict positivity of $\mathbf{b}$, which we cannot easily express in terms of $\mathbf{G}$ and $\mathbf{d}$, and it is not obvious that their proof can be directly used to guarantee convergence. Without guarantees of convergence, we express their method in our nomenclature as Algorithm 2, computing the step as $\mathbf{h}^{(k+1)} \leftarrow -\nabla\phi(\mathbf{x}^{(k)}) \odot \mathbf{x}^{(k)} \oslash (\mathbf{G}\mathbf{x}^{(k)})$. The vague language around selecting the step length τ_k is due to Merritt and Zhang; presumably one can choose it randomly, or always take $\tau_k = (\tau + 1)/2$.

3.1 Other Directions

Imagine, instead, plugging in different values of $\mathbf{b}$ into the $\mathbf{x}^{(k+1)}$ given by Lemma 2.2. We claim that setting $\mathbf{b} = \mathbf{x}^{(k)} + \mathbf{G}^{-1}\mathbf{d}$ would yield the global unconstrained minimizer for $\phi(\mathbf{x})$:

$$\begin{aligned} \mathbf{x}^{(k+1)} &\leftarrow \mathbf{x}^{(k)} - \text{Diag}(\mathbf{x}^{(k)} + \mathbf{G}^{-1}\mathbf{d}) \text{Diag}(\mathbf{G}\mathbf{x}^{(k)} + \mathbf{d})^{-1} (\mathbf{d} + \mathbf{G}\mathbf{x}^{(k)}), \\ &= \mathbf{x}^{(k)} - \text{Diag}(\mathbf{x}^{(k)} + \mathbf{G}^{-1}\mathbf{d}) \mathbf{1}, \\ &= -\mathbf{G}^{-1}\mathbf{d}. \end{aligned} \tag{25}$$

This would seem to be a totally useless computation because we cannot efficiently compute $-\mathbf{G}^{-1}\mathbf{d}$, and it likely violates the positivity constraint. However, this suggests an alternative iterative step that might have accelerated convergence. For example, it might be the case that setting $\mathbf{b} = (\mathbf{x}^{(k)} + \text{Diag}(\text{diag}(\mathbf{G}))^{-1}\mathbf{d})^+$, might give quicker convergence to a solution, where x^+ is the non-negative part of x .

We can also imagine an iterative update based on different descent directions altogether. For example steepest descent, where $\mathbf{h}^{(k+1)} = -\nabla\phi(\mathbf{x}^{(k)})$. Both of these and the method above can be couched as the general iterative method of Algorithm 2. We note, however, that steepest descent will fail to make progress in this formulation when an element of $\mathbf{x}^{(k)}$ is zero and the corresponding element of the gradient is positive. In this case, the step length would

be computed as zero and the algorithm would terminate early. The Lee and Seung step naturally avoids this problem by scaling the step direction proportional element-wise to $\mathbf{x}^{(k)}$. It is not clear, however, whether the denominator part of the Lee and Seung step is necessary, and perhaps taking steps in the direction of $\mathbf{h}^{(k+1)} = -\nabla\phi(\mathbf{x}^{(k)}) \odot \mathbf{x}^{(k)}$ sufficiently good convergence.

Algorithm 2 The General Iterative Quadrative Programming Method

```

function GIQPM_STEP( $\mathbf{x}^{(k)}$ ,  $\mathbf{G}$ ,  $\mathbf{d}$ ,  $\tau \in (0, 1)$ )
  Compute the gradient at  $\mathbf{x}^{(k)}$ :  $\nabla\phi(\mathbf{x}^{(k)}) \leftarrow \mathbf{G}\mathbf{x}^{(k)} + \mathbf{d}$ .
  Somehow choose step  $\mathbf{h}^{(k+1)}$  such that  $\mathbf{h}^{(k+1)\top} \nabla\phi(\mathbf{x}^{(k)}) < 0$ .
  Compute maximum allowable step length:
     $\hat{\alpha}_k \leftarrow \min \left\{ \alpha : \mathbf{x}^{(k)} + \alpha \mathbf{h}^{(k+1)} \geq 0 \right\}$ .
  Compute optimal step length:
     $\alpha_k^* \leftarrow -\nabla\phi(\mathbf{x}^{(k)})^\top \mathbf{h}^{(k+1)} / (\mathbf{h}^{(k+1)\top} \mathbf{G} \mathbf{h}^{(k+1)})$ .
  Choose  $\tau_k \in [\tau, 1)$  and let  $\alpha_k$  be the minimum of  $\tau_k \hat{\alpha}_k$  and  $\alpha_k^*$ .
  Let  $\mathbf{x}^{(k+1)} \leftarrow \mathbf{x}^{(k)} + \alpha_k \mathbf{h}^{(k+1)}$ .
  return  $\mathbf{x}^{(k+1)}$ 
end function
function GIQPM( $\mathbf{G}$ ,  $\mathbf{d}$ ,  $\tau \in (0, 1)$ )
  Initialize  $\mathbf{x}^{(0)} > 0$  and set  $k \leftarrow 0$ .
  while not converged do
    Let  $\mathbf{x}^{(k+1)} \leftarrow \text{GIQPM\_STEP}(\mathbf{x}^{(k)}, \mathbf{G}, \mathbf{d}, \tau)$ .
    Increment  $k \leftarrow k + 1$ .
    Check convergence.
  end while
  return  $\mathbf{x}^{(k)}$ 
end function

```

Converting this to the NMF problem as formulated in Section 2.1.1 is a straightforward, though somewhat tedious, exercise. We give the algorithm in Algorithm 3 for the case where the weighting matrices satisfy the conditions of Equations 4. We note that the restrictions on $\mathbf{W}_{1,R}$ and $\mathbf{W}_{1,L}$ that caused so much headache above are swept under the rug here with the check for maximum allowable step length, $\hat{\alpha}_k$, in GIQPM_STEP.

It is not clear *a priori* whether this algorithm converges quickly. It is very unlikely that Algorithm 2 is competitive for the constrained optimization problem on vectors given in 10. Converting to steepest descent is not recommended due to slow convergence, as noted above. One could swap out the iterative updates of $\mathbf{L}^{(k)}$ and $\mathbf{R}^{(k)}$ in Algorithm 3 with something like a constrained conjugate gradient. [3, 5] However, one should take several steps of conjugate gradient for each k as the optimization problem changes as we update the $\mathbf{L}^{(k)}$ and $\mathbf{R}^{(k)}$.

4 Simulations

The multiplicative and additive algorithms are implemented in the R package RNNMF written by the author. [7] Here we briefly summarize the results of using the code on generated data. We first generate a $\mathbf{Y}$ that exactly equals $\mathbf{L}\mathbf{R}$ for

Algorithm 3 Additive Update Regularized Non-Negative Matrix Factorization, II.

function GIQPM_STEP($\mathbf{X}^{(k)}, \nabla\phi(\mathbf{X}^{(k)}), \mathbf{H}^{(k+1)}, \mathbf{K}^{(k+1)}, \tau_k \in (0, 1)$)

 Compute maximum allowable step length:

$\hat{\alpha}_k \leftarrow \min \{ \alpha : \mathbf{X}^{(k)} + \alpha \mathbf{H}^{(k+1)} \geq 0 \}.$

 Compute optimal step length:

$\alpha_k^* \leftarrow -\text{tr} \left(\nabla\phi(\mathbf{X}^{(k)})^\top \mathbf{H}^{(k+1)} \right) / \text{tr} \left(\mathbf{H}^{(k+1)\top} \mathbf{K}^{(k+1)} \right).$

 (compute these traces without performing matrix multiplies.)

 Let α_k be the minimum of $\tau_k \hat{\alpha}_k$ and $\alpha_k^*.$

 Let $\mathbf{X}^{(k+1)} \leftarrow \mathbf{X}^{(k)} + \alpha_k \mathbf{H}^{(k+1)}.$

return $\mathbf{X}^{(k+1)}$

end function

function PICK_DIRECTION($\mathbf{X}^{(k)}, \nabla\phi(\mathbf{X}^{(k)}), \mathbf{F}$)

 Initialize $\mathbf{H}^{(k)} \leftarrow -\nabla\phi(\mathbf{X}^{(k)}) \odot \mathbf{X}^{(k)} \oslash \mathbf{F}.$

 Elements of $\mathbf{H}^{(k)}$ for which $\mathbf{F}$ and $\mathbf{X}^{(k)}$ are both zero set to $\max(-\nabla\phi(\mathbf{X}^{(k)}), 0).$

 Elements of $\mathbf{H}^{(k)}$ for which $\mathbf{F}$ is zero but $\mathbf{X}^{(k)} > 0$ set to $-\nabla\phi(\mathbf{X}^{(k)}) \odot \mathbf{X}^{(k)}.$

return $\mathbf{H}^{(k+1)}$

end function

function AURNMF($\mathbf{Y}, d, \mathbf{W}_{0,R}, \mathbf{W}_{0,C}, \lambda_{1,L}, \lambda_{1,R}, \lambda_{2,L}, \lambda_{2,R}, \gamma_{2,L}, \gamma_{2,R}, \tau \in (0, 1)$)

 Initialize random $\mathbf{L}^{(0)} > 0$ with d columns and $\mathbf{R}^{(0)} > 0$ with d rows.

 Let $k \leftarrow 0.$

while not converged **do**

 Choose $\tau_k \in [\tau, 1).$

$\mathbf{D} \leftarrow \left(\lambda_{1,L} \mathbf{1} - \mathbf{W}_{0,R} \mathbf{Y} \mathbf{W}_{0,C} \mathbf{R}^{(k)\top} \right).$

$\mathbf{F} \leftarrow \mathbf{W}_{0,R} \mathbf{L}^{(k)} \mathbf{R}^{(k)} \mathbf{W}_{0,C} \mathbf{R}^{(k)\top} + \lambda_{2,L} \mathbf{L}^{(k)} + \gamma_{2,L} \mathbf{L}^{(k)} (\mathbf{1}\mathbf{1}^\top - \mathbf{I})$

$\nabla\phi(\mathbf{L}^{(k)}) \leftarrow \mathbf{F} + \mathbf{D}.$

$\mathbf{H}^{(k+1)} \leftarrow \text{PICK_DIRECTION}(\mathbf{L}^{(k)}, \nabla\phi(\mathbf{L}^{(k)}), \mathbf{F}).$

$\mathbf{K}^{(k+1)} \leftarrow \mathbf{W}_{0,R} \mathbf{H}^{(k+1)} \mathbf{R}^{(k)} \mathbf{W}_{0,C} \mathbf{R}^{(k)\top} + \lambda_{2,L} \mathbf{H}^{(k+1)}.$

 Let $\mathbf{L}^{(k+1)} \leftarrow \text{GIQPM_STEP}(\mathbf{L}^{(k)}, \nabla\phi(\mathbf{L}^{(k)}), \mathbf{H}^{(k+1)}, \mathbf{K}^{(k+1)}, \tau_k).$

$\mathbf{D} \leftarrow \left(\lambda_{1,R} \mathbf{1} - \mathbf{L}^{(k+1)\top} \mathbf{W}_{0,R} \mathbf{Y} \mathbf{W}_{0,C} \right).$

$\mathbf{F} \leftarrow \mathbf{L}^{(k+1)\top} \mathbf{W}_{0,R} \mathbf{L}^{(k+1)} \mathbf{R}^{(k)} \mathbf{W}_{0,C} + \lambda_{2,R} \mathbf{R}^{(k)} + \gamma_{2,R} \mathbf{R}^{(k)} (\mathbf{1}\mathbf{1}^\top - \mathbf{I})$

$\nabla\phi(\mathbf{R}^{(k)}) \leftarrow \mathbf{F} + \mathbf{D}.$

$\mathbf{H}^{(k+1)} \leftarrow \text{PICK_DIRECTION}(\mathbf{R}^{(k)}, \nabla\phi(\mathbf{R}^{(k)}), \mathbf{F}).$

$\mathbf{K}^{(k+1)} \leftarrow \mathbf{L}^{(k+1)\top} \mathbf{W}_{0,R} \mathbf{L}^{(k+1)} \mathbf{H}^{(k+1)} \mathbf{W}_{0,C} + \lambda_{2,R} \mathbf{H}^{(k+1)}$

 Let $\mathbf{R}^{(k+1)} \leftarrow \text{GIQPM_STEP}(\mathbf{R}^{(k)}, \nabla\phi(\mathbf{R}^{(k)}), \mathbf{H}^{(k+1)}, \mathbf{K}^{(k+1)}, \tau_k).$

 Increment $k \leftarrow k + 1.$

 Check convergence.

end while

return $\mathbf{x}^{(k)}$

end function

some non-negative uniformly random L and R . In the first simulations we take L to be 30×2 , and R to be 2×8 . We randomly generate starting iterates, 30×2 matrix $L^{(0)}$ and 2×8 matrix $R^{(0)}$. We use the same starting point for both the additive and multiplicative algorithms. We compute the Frobenius norm of the error, which is to say

$$\sqrt{\text{tr} \left((Y - L^{(k)} R^{(k)})^\top (Y - L^{(k)} R^{(k)}) \right)},$$

for each iteration.

In Figure 1 we plot the error versus step for the additive and multiplicative methods. The additive method uses the optimal step size in the chosen direction. While not shown here, choosing the naive step without optimizing converges slower in iterations but may ultimately be faster computationally. The additive method outperforms the multiplicative method in convergence per step. Convergence of both is “lumpy”, with apparent phase changes where, we suspect, different directions become the dominant source of approximation error, and then are smoothed out.

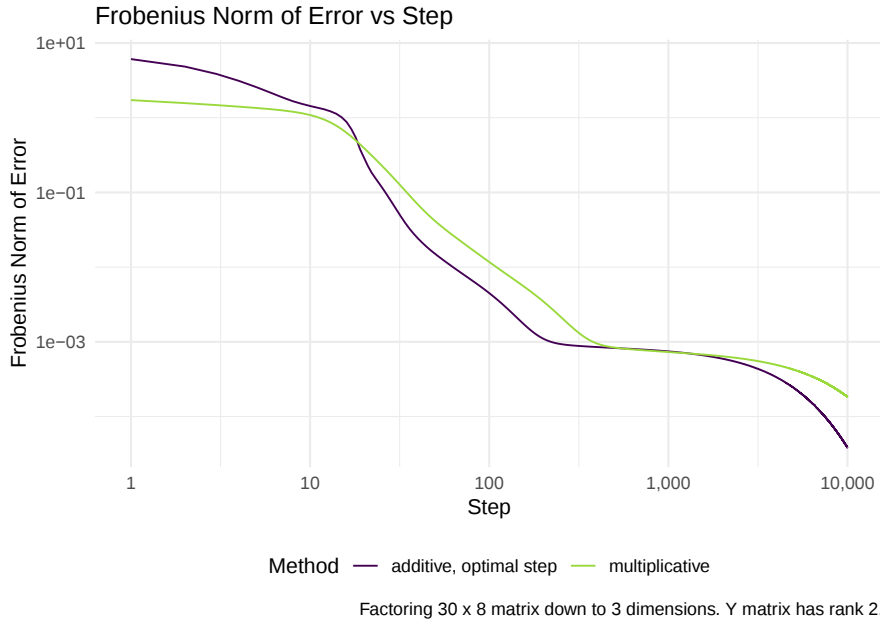


Figure 1: The Frobenius norm is plotted versus step for two methods for a small problem.

We then apply the code to another test case. Here we generate L to be 40×3 , and R to be 3×10 . Again we test both the additive and multiplicative algorithms, but test the effect of sparse starting iterates. We randomly generate starting iterates, 40×4 matrix $L^{(0)}$ and 4×10 matrix $R^{(0)}$. First we do this where about $\frac{1}{3}$ of the elements of $L^{(0)}$ and $R^{(0)}$ are zero, testing both algorithms. We repeat the experiment for the same Y , but generate strictly positive $L^{(0)}$ and $R^{(0)}$.

In Figure 2 we plot the error versus step for both methods and both choices of starting iterates. Again we see somewhat slower convergence for the multiplicative method, at least measured in iterates. Moreover, the multiplicative method fails to converge for the sparse starting iterates. As we noted above, the multiplicative method cannot change an element from zero to non-negative, thus the lack of convergence for sparse initial estimates is not surprising.

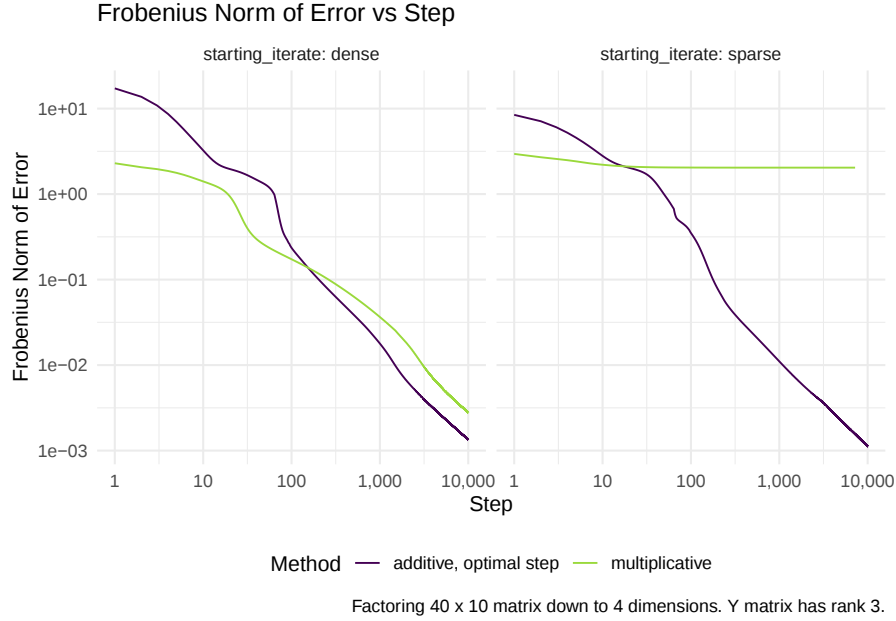


Figure 2: The Frobenius norm is plotted versus step for two methods for a small problem. Starting iterates are taken to be sparse or dense.

References

- [1] Edward F. Gonzalez and Yin Zhang. Accelerating the Lee-Seung algorithm for nonnegative matrix factorization. Technical Report TR05-02, Rice University, 2005. URL <https://repository.rice.edu/bitstream/handle/1911/102034/TR05-02.pdf>.
- [2] Daniel D. Lee and H. Sebastian Seung. Algorithms for non-negative matrix factorization. In T. K. Leen, T. G. Dietterich, and V. Tresp, editors, *Advances in Neural Information Processing Systems 13*, pages 556–562. MIT Press, 2001. URL <http://papers.nips.cc/paper/1861-algorithms-for-non-negative-matrix-factorization.pdf>.
- [3] Can Li. A conjugate gradient type method for the nonnegative constraints optimization problems. *Journal of Applied Mathematics*, 2013:986317, 2013. doi: 10.1155/2013/986317. URL <https://doi.org/10.1155/2013/986317>.
- [4] Michael Merritt and Yin Zhang. Interior-point gradient method for large-scale totally nonnegative least squares problems. *Journal of optimization*

- theory and applications*, 126(1):191–202, 2005. URL <https://repository.rice.edu/bitstream/handle/1911/102020/TR04-08.pdf>.
- [5] J. Nocedal and S. J. Wright. *Numerical Optimization*. Springer series in operations research and financial engineering. Springer, 2006. ISBN 9780387400655. URL <http://books.google.com/books?id=VbHYoSyelFcC>.
- [6] Steven E. Pav. System and method for unmixing spectroscopic observations with nonnegative matrix factorization, 2012. URL <https://patentscope.wipo.int/search/en/detail.jsf?docId=US42758160>.
- [7] Steven E. Pav. *rnnmf: Regularized Non-negative Matrix Factorization*, 2024. URL <https://github.com/shabbychef/rnnmf>. R package version 0.3.0.
- [8] Farid Saberi-Movahed, Kamal Berahman, Razieh Sheikhpour, Yuefeng Li, and Shirui Pan. Nonnegative matrix factorization in dimensionality reduction: A survey, 2024. URL <https://arxiv.org/abs/2405.03615>.
- [9] Kathrin Schäcke. On the Kronecker product, 2013. URL <https://www.math.uwaterloo.ca/~hwolkowi/henry/reports/kronthesisschaecke04.pdf>.

A Matrix Identities

The following identities are useful for switching between matrix and vectorized representations. [9]

$$\text{vec}(\text{BCD}) = (\text{D}^\top \otimes \text{B}) \text{vec}(\text{C}). \quad (26)$$

$$\text{tr}(\text{A}^\top \text{C}) = \text{vec}(\text{A})^\top \text{vec}(\text{C}). \quad (27)$$

$$\text{tr}(\text{A}^\top \text{BCD}) = \text{vec}(\text{A})^\top (\text{D}^\top \otimes \text{B}) \text{vec}(\text{C}). \quad (28)$$

$$\text{tr}(\text{A}^\top \text{BAD}) = \|\text{vec}(\text{A})\|_{\text{D}^\top \otimes \text{B}}^2. \quad (29)$$

$$(\text{A} \otimes \text{B})^\top = \text{A}^\top \otimes \text{B}^\top. \quad (30)$$