

Convergence of Gradient Descent with Small Initialization for Unregularized Matrix Completion

Jianhao Ma

Joint work with Salar Fattahi

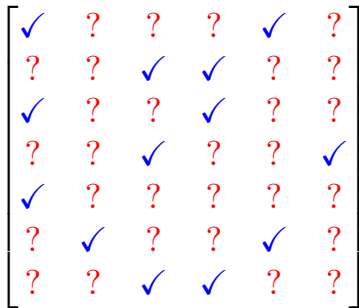
University of Michigan

COLT 2024

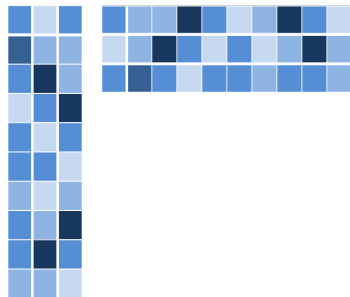
Matrix completion

Goal

Recover a low-rank matrix from a subset of its entries.



(a) matrix completion



(b) matrix factorization¹

¹Adapted from https://yuxinchen2020.github.io/slides/NoisyMC_slides.pdf

Mathematical formulation

Unregularized nonconvex optimization

$$\min_{\mathbf{U} \in \mathbb{R}^{d \times r'}} f(\mathbf{U}) = \frac{1}{4p} \left\| \mathcal{P}_{\Omega}(\mathbf{X}^* - \mathbf{U}\mathbf{U}^{\top}) \right\|_{\text{F}}^2 = \frac{1}{4p} \sum_{(i,j) \in \Omega} \left(X_{i,j}^* - [\mathbf{U}\mathbf{U}^{\top}]_{i,j} \right)^2.$$

Mathematical formulation

Unregularized nonconvex optimization

$$\min_{U \in \mathbb{R}^{d \times r'}} f(U) = \frac{1}{4p} \left\| \mathcal{P}_{\Omega}(\mathbf{X}^* - UU^{\top}) \right\|_{\text{F}}^2 = \frac{1}{4p} \sum_{(i,j) \in \Omega} \left(X_{i,j}^* - [UU^{\top}]_{i,j} \right)^2.$$

- $\mathbf{X}^*$ is a PSD matrix with rank $r \ll d$.

Mathematical formulation

Unregularized nonconvex optimization

$$\min_{U \in \mathbb{R}^{d \times r'}} f(U) = \frac{1}{4p} \left\| \mathcal{P}_{\Omega}(\mathbf{X}^* - \mathbf{U}\mathbf{U}^{\top}) \right\|_{\text{F}}^2 = \frac{1}{4p} \sum_{(i,j) \in \Omega} \left(X_{i,j}^* - [\mathbf{U}\mathbf{U}^{\top}]_{i,j} \right)^2.$$

- $\mathbf{X}^*$ is a PSD matrix with rank $r \ll d$.
- **Search rank:** Overparameterization ($r' > r$) v.s. exact-parameterization ($r' = r$).

Mathematical formulation

Unregularized nonconvex optimization

$$\min_{U \in \mathbb{R}^{d \times r'}} f(U) = \frac{1}{4p} \left\| \mathcal{P}_\Omega(\mathbf{X}^\star - \mathbf{U}\mathbf{U}^\top) \right\|_F^2 = \frac{1}{4p} \sum_{(i,j) \in \Omega} \left(X_{i,j}^\star - [\mathbf{U}\mathbf{U}^\top]_{i,j} \right)^2.$$

- $\mathbf{X}^\star$ is a PSD matrix with rank $r \ll d$.
- **Search rank:** Overparameterization ($r' > r$) v.s. exact-parameterization ($r' = r$).
- **Bernoulli sampling model:** $\mathbf{1}_{\{(i,j) \in \Omega\}} \stackrel{i.i.d.}{\sim} \text{Ber}(p)$.

Incoherence

- **Incoherence:** $\mathbf{X}^* = \mathbf{V}^* \mathbf{\Sigma}^* \mathbf{V}^{*\top}$ satisfies $\|\mathbf{V}^*\|_{2,\infty} \leq \sqrt{\frac{\mu r}{d}}$ where $\mu \geq 1$ and $\|\mathbf{A}\|_{2,\infty} := \max_i \{\|\mathbf{A}_{i,\cdot}\|\}$.

Incoherence

- **Incoherence:** $\mathbf{X}^* = \mathbf{V}^* \mathbf{\Sigma}^* \mathbf{V}^{*\top}$ satisfies $\|\mathbf{V}^*\|_{2,\infty} \leq \sqrt{\frac{\mu r}{d}}$ where $\mu \geq 1$ and $\|\mathbf{A}\|_{2,\infty} := \max_i \{\|\mathbf{A}_{i,\cdot}\|\}$.
- Incoherence determines the **difficulty** (sample complexity) of matrix completion.

$$\underbrace{\begin{bmatrix} 1 & 0 & 0 & \cdots & 0 \\ 0 & 0 & 0 & \cdots & 0 \\ \vdots & \vdots & \vdots & & \\ 0 & 0 & 0 & \cdots & 0 \end{bmatrix}}_{\text{hard } \mu=d} \quad \text{v.s.} \quad \underbrace{\begin{bmatrix} 1 & 1 & 1 & \cdots & 1 \\ 1 & 1 & 1 & \cdots & 1 \\ \vdots & \vdots & \vdots & & \\ 1 & 1 & 1 & \cdots & 1 \end{bmatrix}}_{\text{easy } \mu=1}$$

Prior arts

- **Gradient descent (GD):** $U_{t+1} = U_t - \eta \nabla f(U_t)$.

Prior arts

- **Gradient descent (GD):** $U_{t+1} = U_t - \eta \nabla f(U_t)$.
- GD and its variants converge to the global minimum efficiently for matrix sensing and other similar problems in both exactly and over-parameterized settings [CLC19].

Prior arts

- **Gradient descent (GD):** $U_{t+1} = U_t - \eta \nabla f(U_t)$.
- GD and its variants converge to the global minimum efficiently for matrix sensing and other similar problems in both exactly and over-parameterized settings [CLC19].
- Their analysis relies on the restricted isometry property (RIP).

Prior arts

- **Gradient descent (GD):** $U_{t+1} = U_t - \eta \nabla f(U_t)$.
- GD and its variants converge to the global minimum efficiently for matrix sensing and other similar problems in both exactly and over-parameterized settings [CLC19].
- Their analysis relies on the restricted isometry property (RIP).
- Matrix completion only satisfies RIP in the incoherent region $\left\{U : \|U\|_{2,\infty} \lesssim \sqrt{\frac{\mu r}{d}}\right\}$.

Prior arts (cont.)

- **Explicit regularization:** add extra regularizer [GJZ17] or projection [ZL16].

Prior arts (cont.)

- **Explicit regularization:** add extra regularizer [GJZ17] or projection [ZL16].
- **Implicit regularization:**
 - **Local result [MWCC18]:** Suppose that ① $r' = r$; ② U_0 is close to $U^* = X^{*1/2}$. Then U_t converges to U^* at a linear rate.
 - **Global result [KC22]:** Suppose that ① $r' = r = 1$; ② $U_0 \approx 0$. Then U_t converges to U^* at a linear rate.

Prior arts (cont.)

- **Explicit regularization:** add extra regularizer [GJZ17] or projection [ZL16].
- **Implicit regularization:**
 - **Local result [MWCC18]:** Suppose that ① $r' = r$; ② U_0 is close to $U^* = X^{*1/2}$. Then U_t converges to U^* at a linear rate.
 - **Global result [KC22]:** Suppose that ① $r' = r = 1$; ② $U_0 \approx 0$. Then U_t converges to U^* at a linear rate.

Question: Does GD with small initialization converge to the true solution in the **general rank- r** and **overparameterized** setting?

Overview of our results

Algorithm	Sample complexity	Computational complexity	Global	Exact	Over-param.
[MWCC18]	$dr^3 \log^3(d)$	$\kappa^2 \log\left(\frac{1}{\epsilon}\right)$	✗	✓	✗
[CLL20]	$dr^2 \log(d)$	$\kappa^2 \log\left(\frac{1}{\epsilon}\right)$	✗	✓	✗
[KC22]	$d \log^{22}(d)$	$\log\left(\frac{1}{\epsilon}\right)$	✓	*	✗
Ours (Thm 1)	$dr^9 \log^8(d)$	$\kappa^4 \log\left(\frac{1}{\epsilon}\right)$	✓	✓	✗
Ours (Thm 2)	$dr^9 \log^2(d) \log^6\left(\frac{1}{\epsilon}\right)$	$\kappa^4 \log^4\left(\frac{1}{\epsilon}\right)$	✓	✓	✓

- * This result only holds for $r' = r = 1$.
- κ is the condition number.

Exactly-parameterized regime

$$\boxed{U_{t+1} = U_t - \eta \nabla f(U_t), \quad \text{where} \quad U_0 = \alpha B} \quad (\text{GD})$$

Theorem 1 (Theorem 3 in [MF24])

Suppose that the sampling rate $p \gtrsim \frac{\kappa^6 \mu^4 r^9 \log^8(d)}{d}$. Consider GD with the step-size $\eta \asymp \frac{\mu r}{\sqrt{pd}\sigma_1^*}$ and the initial point $U_0 = \alpha B$, where $\alpha \asymp \frac{\sigma_r^*}{\kappa^{1.5}d}$. Given any accuracy $\epsilon > 0$, with probability at least $1 - O\left(\frac{1}{d^3}\right)$ and after $T \lesssim \frac{1}{\eta\sigma_r^*} \log\left(\frac{1}{\epsilon}\right)$ iterations, we have

$$\left\| U_T U_T^\top - X^* \right\|_F \leq \epsilon.$$

Overparameterized regime

Theorem 2 (Theorem 2 in [MF24])

Suppose that the sampling rate $p \gtrsim \frac{\kappa^6 \mu^4 r^9 \log^6(\frac{1}{\alpha}) \log^2(d)}{d}$. Consider GD with the step-size $\eta \asymp \frac{\mu r}{\sqrt{pd\sigma_1^*}}$ and the initial point $\mathbf{U}_0 = \alpha \mathbf{B}$, where $0 < \alpha \leq \sqrt{\frac{\sigma_1^*}{d}}$. With probability at least $1 - \frac{2}{d^3}$, after $T \lesssim \frac{1}{\eta \sigma_r^*} \log\left(\frac{1}{\alpha}\right)$ iterations, we have

$$\left\| \mathbf{U}_T \mathbf{U}_T^\top - \mathbf{X}^* \right\|_{\text{F}} \lesssim \sqrt{\frac{\sigma_1^* \kappa^2 \mu r^2}{p}} \alpha.$$

Remark: To achieve ϵ -accuracy, it suffices to set $\alpha \lesssim \sqrt{\frac{p}{\sigma_1^* \kappa^2 \mu r^2}} \epsilon$.

Proof outline

- **Dynamic signal-residual decomposition (similar to [LMZ18]²):** Project each iterate U_t onto a dynamic rank- r subspace V_t and its complement:

$$U_t = \underbrace{\mathcal{P}_{V_t} U_t}_{\text{rank-}r} + \underbrace{\mathcal{P}_{V_t}^\perp U_t}_{\approx 0}.$$

²COLT 2018 best paper award.

Proof outline

- **Dynamic signal-residual decomposition (similar to [LMZ18]²):** Project each iterate U_t onto a dynamic rank- r subspace V_t and its complement:

$$U_t = \underbrace{\mathcal{P}_{V_t} U_t}_{\text{rank-}r} + \underbrace{\mathcal{P}_{V_t^\perp} U_t}_{\approx 0}.$$

- **Weakly-coupled leave-one-out (WC-LOO) analysis:** Show that the implicit regularization effect of the low-rank part $\mathcal{P}_{V_t} U_t$:

$$\|\mathcal{P}_{V_t} U_t\|_{2,\infty} \lesssim \sqrt{\frac{\sigma_1^* \mu r}{d}}, \quad \forall t = 0, \dots, T.$$

²COLT 2018 best paper award.

Classic leave-one-out analysis

- **Idea:** decouple the statistical correlation between algorithm and train set.

Classic leave-one-out analysis

- **Idea:** decouple the statistical correlation between algorithm and train set.
- **Leave-one-out sequences:** For each $1 \leq l \leq d$, $\{U_t^{(l)}\}$ is generated by

$$U_{t+1}^{(l)} = \left(I - \eta \frac{1}{p} \mathcal{P}_{\Omega^{(l)}} \left(U_t^{(l)} U_t^{(l)\top} - X^* \right) \right) U_t^{(l)}.$$

$$\underbrace{\begin{bmatrix} 1 & 0 & 0 & 1 & 0 \\ 0 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 1 & 0 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 1 \end{bmatrix}}_{\text{sampling matrix } \Omega} \implies \underbrace{\begin{bmatrix} 1 & p & 0 & 1 & 0 \\ p & p & p & p & p \\ 0 & p & 0 & 1 & 0 \\ 1 & p & 1 & 0 & 0 \\ 0 & p & 0 & 0 & 1 \end{bmatrix}}_{\text{leave-one-out version } \Omega^{(2)}}$$

Classic leave-one-out analysis

- **Idea:** decouple the statistical correlation between algorithm and train set.
- **Leave-one-out sequences:** For each $1 \leq l \leq d$, $\{U_t^{(l)}\}$ is generated by

$$U_{t+1}^{(l)} = \left(I - \eta \frac{1}{p} \mathcal{P}_{\Omega^{(l)}} \left(U_t^{(l)} U_t^{(l)\top} - X^* \right) \right) U_t^{(l)}.$$

$$\underbrace{\begin{bmatrix} 1 & 0 & 0 & 1 & 0 \\ 0 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 1 & 0 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 1 \end{bmatrix}}_{\text{sampling matrix } \Omega} \implies \underbrace{\begin{bmatrix} 1 & p & 0 & 1 & 0 \\ p & p & p & p & p \\ 0 & p & 0 & 1 & 0 \\ 1 & p & 1 & 0 & 0 \\ 0 & p & 0 & 0 & 1 \end{bmatrix}}_{\text{leave-one-out version } \Omega^{(2)}}$$

Key observation: The l -th row of $U_t^{(l)}$ is *deterministic* for all t !

Control incoherence

$$\|U_t\|_{2,\infty} \leq \underbrace{\max_{1 \leq l \leq d} \left\{ \left\| \left(U^\star - U_t^{(l)} R_t^{(l)} \right)_{l,\cdot} \right\| \right\}}_{\text{leave-one-out error}} + \underbrace{\max_{1 \leq l \leq d} \left\{ \left\| U_t - U_t^{(l)} H_t^{(l)} \right\|_F \right\}}_{\text{proximal error}} + \sqrt{\frac{\sigma_1^\star \mu r}{d}}.$$

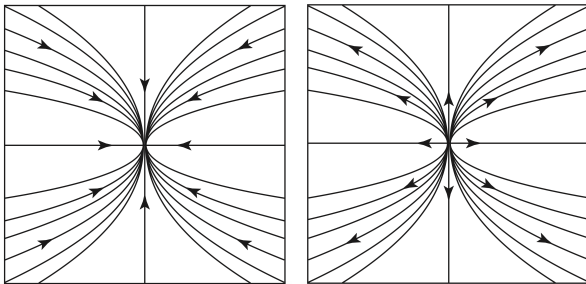
- **Leave-one-out error:** By our construction of $\Omega^{(l)}$, the leave-one-out error is *deterministic*!

Control incoherence

$$\|U_t\|_{2,\infty} \leq \underbrace{\max_{1 \leq l \leq d} \left\{ \left\| \left(U^\star - U_t^{(l)} R_t^{(l)} \right)_{l,\cdot} \right\| \right\}}_{\text{leave-one-out error}} + \underbrace{\max_{1 \leq l \leq d} \left\{ \left\| U_t - U_t^{(l)} H_t^{(l)} \right\|_F \right\}}_{\text{proximal error}} + \sqrt{\frac{\sigma_1^\star \mu r}{d}}.$$

- **Leave-one-out error:** By our construction of $\Omega^{(l)}$, the leave-one-out error is *deterministic*!
- **Proximal error:** Suppose that $U_t \approx U^\star$. Then, due to the local restricted *strong convexity* of the loss function, the proximal error decreases geometrically fast.

Local minimum v.s. saddle point



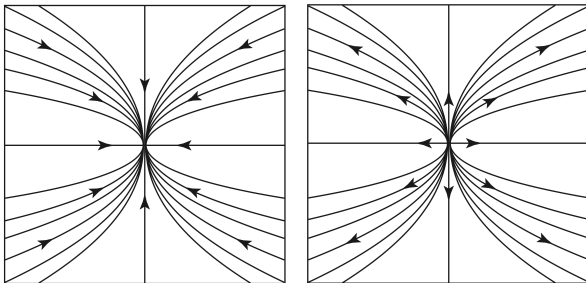
(a) U^* is stable

(b) 0 is unstable

Figure: Gradient flow around U^* and 0 .

Local minimum v.s. saddle point

- U^* is a local minimum: The distance between U_t and $U_t^{(l)}$ decreases linearly.



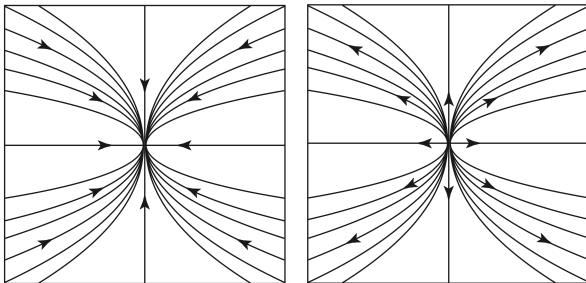
(a) U^* is stable

(b) 0 is unstable

Figure: Gradient flow around U^* and 0 .

Local minimum v.s. saddle point

- U^* is a **local minimum**: The distance between U_t and $U_t^{(l)}$ decreases linearly.
- 0 is a **saddle point**: Under small initialization, the distance between U_t and $U_t^{(l)}$ increases exponentially fast.



(a) U^* is stable

(b) 0 is unstable

Figure: Gradient flow around U^* and 0 .

Weakly-coupled leave-one-out sequences

original sequence

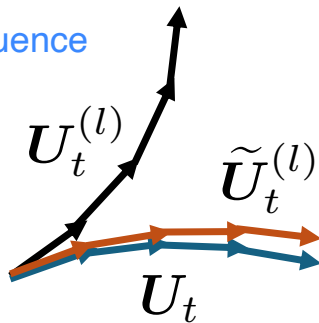
classic LOO sequence

$$U_{t+1} = \left(I - \eta \frac{1}{p} \mathcal{P}_{\Omega} \left(V_t \Sigma_t V_t^{\top} - X^{\star} \right) \right) U_t$$

$$U_{t+1}^{(l)} = \left(I - \eta \frac{1}{p} \mathcal{P}_{\Omega^{(l)}} \left(V_t^{(l)} \Sigma_t^{(l)} V_t^{(l)\top} - X^{\star} \right) \right) U_t^{(l)}$$

$$\tilde{U}_{t+1}^{(l)} = \left(I - \eta \frac{1}{p} \mathcal{P}_{\Omega^{(l)}} \left(\tilde{V}_t^{(l)} \Sigma_t \tilde{V}_t^{(l)\top} - X^{\star} \right) \right) \tilde{U}_t^{(l)}$$

WC-LOO sequence



Key observation: $\tilde{U}_t^{(l)}$ is much closer to U_t than $U_t^{(l)}$ since we replace $\Sigma_t^{(l)}$ with Σ_t .

²Here we define $\Sigma_t = V_t^{\top} U_t U_t^{\top} V_t \in \mathbb{R}^{r \times r}$.

Control incoherence via WC-LOO analysis

$$\|U_t\|_{2,\infty} \leq \underbrace{\max_{1 \leq l \leq d} \left\{ \left\| \left(U^* - \tilde{U}_t^{(l)} \tilde{R}_t^{(l)} \right)_{l,\cdot} \right\| \right\}}_{\text{WC-LOO error}} + \underbrace{\max_{1 \leq l \leq d} \left\{ \left\| U_t - \tilde{U}_t^{(l)} \tilde{H}_t^{(l)} \right\|_F \right\}}_{\text{WC-proximal error}} + \sqrt{\frac{\sigma_1^* \mu r}{d}}$$

- **Good news ☺:** The WC-proximal error grows mildly for all $0 \leq t \leq T$.

Control incoherence via WC-LOO analysis

$$\|U_t\|_{2,\infty} \leq \underbrace{\max_{1 \leq l \leq d} \left\{ \left\| \left(U^* - \tilde{U}_t^{(l)} \tilde{R}_t^{(l)} \right)_{l,\cdot} \right\| \right\}}_{\text{WC-LOO error}} + \underbrace{\max_{1 \leq l \leq d} \left\{ \left\| U_t - \tilde{U}_t^{(l)} \tilde{H}_t^{(l)} \right\|_F \right\}}_{\text{WC-proximal error}} + \sqrt{\frac{\sigma_1^* \mu r}{d}}$$

- **Good news** ☺: The WC-proximal error grows mildly for all $0 \leq t \leq T$.
- **Bad news** ☹: The WC-LOO error is no longer deterministic since we use Σ_t .

Control incoherence via WC-LOO analysis

$$\|U_t\|_{2,\infty} \leq \underbrace{\max_{1 \leq l \leq d} \left\{ \left\| \left(U^* - \tilde{U}_t^{(l)} \tilde{R}_t^{(l)} \right)_{l,\cdot} \right\| \right\}}_{\text{WC-LOO error}} + \underbrace{\max_{1 \leq l \leq d} \left\{ \left\| U_t - \tilde{U}_t^{(l)} \tilde{H}_t^{(l)} \right\|_F \right\}}_{\text{WC-proximal error}} + \sqrt{\frac{\sigma_1^* \mu r}{d}}$$

- **Good news ☺:** The WC-proximal error grows mildly for all $0 \leq t \leq T$.
- **Bad news ☹:** The WC-LOO error is no longer deterministic since we use Σ_t .
- **Good news ☺:** Σ_t is an $r \times r$ small matrix. Hence, we can use a dynamic covering trick to bound the deviation uniformly. The price is to increase the sample complexity by a $\text{poly}(r)$ factor.

Take-home message

- GD with small initialization converges to the ground truth without any explicit regularization.
- It achieves nearly optimal sample complexity and computational complexity.
- We propose a powerful tool called weakly-coupled leave-one-out analysis, which can decouple the statistical correlation between algorithm and training dataset in the **nonconvex** regime.

References



Yuejie Chi, Yue M Lu, and Yuxin Chen.

Nonconvex optimization meets low-rank matrix factorization: An overview.
IEEE Transactions on Signal Processing, 67(20):5239–5269, 2019.



Ji Chen, Dekai Liu, and Xiaodong Li.

Nonconvex rectangular matrix completion via gradient descent without $\ell_{2,\infty}$ regularization.
IEEE Transactions on Information Theory, 66(9):5806–5841, 2020.



Rong Ge, Chi Jin, and Yi Zheng.

No spurious local minima in nonconvex low rank problems: A unified geometric analysis.
In *International Conference on Machine Learning*, pages 1233–1242. PMLR, 2017.



Daesung Kim and Hye Won Chung.

Rank-1 matrix completion with gradient descent and small random initialization.
arXiv preprint arXiv:2212.09396, 2022.