



優化筆記

[Uwni](#) satanyalin@gmail.com

目录

1	Gradient Descent	2
1.1	Equivalent Problems	2
1.2	Unconstrained Case	3
1.3	Constrained Case	5
1.3.1	Projected Gradient Method	6
1.3.2	Example: Linear Constraint	6
1.3.3	Penalty Method	8
1.4	Convergence	8
1.5	Convex Optimization	9
2	光滑流形上的優化	10
2.1	幺正流形	10
2.2	Riemannian 梯度	10
2.3	回撤	10

1 GRADIENT DESCENT

Optimization problems can be divided into minimization and maximization categories, and maximization problems can always be transformed into equivalent minimization problems. Therefore, in the following text, we will focus on minimization problems. Gradient descent is an iterative optimization algorithm used to find the local minimum of a function. It gradually approaches the minimum point by moving in the opposite direction of the function's gradient.

1.1 EQUIVALENT PROBLEMS

First, let us prove some lemmas about monotonicity.

PROPOSITION 1 Order Preservation of Strictly Monotonic Functions *If f is a strictly increasing function, then*

$$x_1 < x_2 \Leftrightarrow f(x_1) < f(x_2)$$

That is, when the output strictly increases, the input must strictly increase.

Proof. ($\rightarrow$) This is the definition of strictly increasing, $x_1 < x_2 \rightarrow f(x_1) < f(x_2)$

($\leftarrow$) By contradiction, if $x_1 \geq x_2$ then $f(x_1) \geq f(x_2)$, which is a contradiction. ■

PROPOSITION 2 Strictly Monotonic Function $\rightarrow$ Injection *If f is a strictly monotonic function, then f is injective. That is,*

$$x_1 = x_2 \Leftrightarrow f(x_1) = f(x_2)$$

Proof. ($\rightarrow$) This follows from the definition of functions.

($\leftarrow$) We need to prove injectivity. Taking f strictly increasing on X as an example. Let $x_1, x_2 \in X, f(x_1) = f(x_2)$. Suppose $x_1 < x_2$, then by strict monotonicity, $x_1 < x_2 \rightarrow f(x_1) < f(x_2)$, which is a contradiction. Similarly $x_1 \not> x_2$, therefore $x_1 = x_2$. ■

PROPOSITION 3 Strictly Monotonic Functions Preserve Extreme Points *Let $Y \subseteq \mathbf{R}$, $f : X \rightarrow Y$ be an arbitrary function, and $g : Y \rightarrow \mathbf{R}$ be a strictly increasing function. Then $g \circ f$ and f have the same extreme points. Conversely, they have opposite extreme points.*

Proof. By the lemma, we know $x_1 \leq x_2 \Leftrightarrow g(x_1) \leq g(x_2)$ Thus for some point $x^* \in X, \exists \delta > 0, \forall x \in U(x^*, \delta)$

$$f(x^*) \leq f(x) \Leftrightarrow g(f(x^*)) \leq g(f(x))$$

This proves that a minimum point of f is also a minimum point of $g \circ f$, and a minimum point of $g \circ f$ is also a minimum point of f . ■

Based on this, for the optimization problem

$$\arg \min f(x)$$

it always has the same solution as $\arg \min g \circ f(x)$, if g is a strictly increasing function. Or $\arg \max g \circ f(x)$, if g is a strictly decreasing function.

□ EXAMPLE 1

Consider an interesting matrix optimization problem that demonstrates the equivalence of different objective functions under monotonic transformations.

Let $X = (x_{ij})_{n \times n}$ be a non-negative matrix with column sum constraints: $\sum_{i=1}^n x_{ij} = c$ for all j (where c is a constant).

Define two objective functions:

$$f_1(X) = \frac{\sum_{i=1}^n x_{ii}}{\sum_{i \neq j} x_{ij}}$$

(diagonal elements / off-diagonal elements)

$$f_2(X) = \frac{\sum_{i=1}^n x_{ii}}{\sum_{i,j=1}^n x_{ij}}$$

(diagonal elements / all elements)

Due to the column sum constraint, the total sum of all elements is: $\sum_{i,j=1}^n x_{ij} = nc$

Therefore: $f_2(X) = (\sum_{i=1}^n x_{ii}) / (nc)$

The sum of off-diagonal elements is: $\sum_{i \neq j} x_{ij} = nc - \sum_{i=1}^n x_{ii}$

So: $f_1(X) = (\sum_{i=1}^n x_{ii}) / (nc - \sum_{i=1}^n x_{ii})$

Key Observation: Let $s = \sum_{i=1}^n x_{ii}$, then:

- $f_2 = s / (nc)$
- $f_1 = s / (nc - s)$

These two functions have a monotonic relationship: $f_1 = f_2 / (1 - f_2) \cdot 1/n$

Since f_1 is a strictly increasing function of f_2 (in the range $f_2 < 1$), we have:

The optimization problems $\max f_1(X)$ and $\max f_2(X)$ are equivalent and have the same optimal solution.

□

1.2 UNCONSTRAINED CASE

$\mathbf{R}^n$ is the n -dimensional Euclidean space. $\mathbf{x} \in \mathbf{R}^n$, $f : \mathbf{R}^n \rightarrow \mathbf{R}$ is a differentiable function. Finding the minimum point and minimum value of f is the optimization problem

$$\min_{\mathbf{x} \in \mathbf{R}^n} f(\mathbf{x})$$

The iterative equation is

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \alpha \nabla f(\mathbf{x}_k)$$

where the step size α is a positive number that determines the update magnitude at each iteration. $\nabla f(\mathbf{x}_k)$ is the gradient of function f at point $\mathbf{x}_k$. The algorithm's goal is to make

$\mathbf{x}_k \rightarrow \mathbf{x}^*$ as $k \rightarrow \infty$ where $\mathbf{x}^* \in \arg \min f(\mathbf{x})$. that is to say, $\mathbf{x}^*$ is a minimum. The pseudocode is as follows:

Input: initial point $\mathbf{x}_0$, step length $\alpha > 0$, tolerance $\varepsilon > 0$, max iterations N
Output: $\mathbf{x}^*$

```

1  $k \leftarrow 0$ 
2 while  $k < N$ 
3    $\mathbf{g}_k \leftarrow \nabla f(\mathbf{x}_k)$ 
4    $\mathbf{x}_{k+1} \leftarrow \mathbf{x}_k - \alpha \mathbf{g}_k$ 
5   if  $f(\mathbf{x}_{k+1}) < \varepsilon$  then
6      $\perp$  return  $\mathbf{x}_{k+1}$ 
7    $k \leftarrow k + 1$ 
8 return  $\mathbf{x}_k$ 

```

Algorithm 1 Pseudocode of Gradient Descent

Let us look at an example

$$\min_{(x_1, x_2) \in \mathbf{R}^2} f(x_1, x_2)$$

where $f(x_1, x_2) = x_1^2 + x_1 x_2 + x_2^2$. First, we know through analytical methods that for $x_1, x_2 \in \mathbf{R}$

$$x_1^2 + x_1 x_2 + x_2^2 = (x_1, x_2) \begin{bmatrix} 1 & \frac{1}{2} \\ \frac{1}{2} & 1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} \geq 0$$

Equality holds if and only if $x_1 = 0, x_2 = 0$. Therefore, the minimum value 0 is achieved at the origin. Next, we use gradient descent to solve this.

$$\nabla f \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} 2x_1 + x_2 \\ 2x_2 + x_1 \end{bmatrix}$$

$$\begin{bmatrix} x_1 \\ x_2 \end{bmatrix}_{k+1} = \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}_k - \alpha \begin{bmatrix} 2x_1 + x_2 \\ 2x_2 + x_1 \end{bmatrix}_k$$

We set the initial condition as $\mathbf{x}_0 = (1.0, 2.0)^T$, step size $\alpha = 0.1$, and stopping condition as $f \leq 10^{-20}$. After 212 iterations, the function value reaches 9.925765507684842e-21. The trajectory left by each iteration in the feasible region is shown in the figure below. The black solid lines in the figure are the contour lines of function f , and the arrows indicate the gradient field. The coloring indicates the magnitude of the function value, with darker colors representing larger function values. The red points represent the positions updated at each iteration, and the connecting lines are the iteration trajectories. It can be seen that each iteration moves opposite to the gradient direction with step size proportional to the gradient magnitude, and the iteration points gradually approach the origin—the theoretical minimum point.

When we increase the step size to 0.4, after 44 iterations, the function value reaches 7.503260807194337e-21.

When we increase the step size to 0.5, after 35 iterations, the function value reaches 5.929230630780102e-21.

If the step size is increased to 0.6, it leads to divergence. Therefore

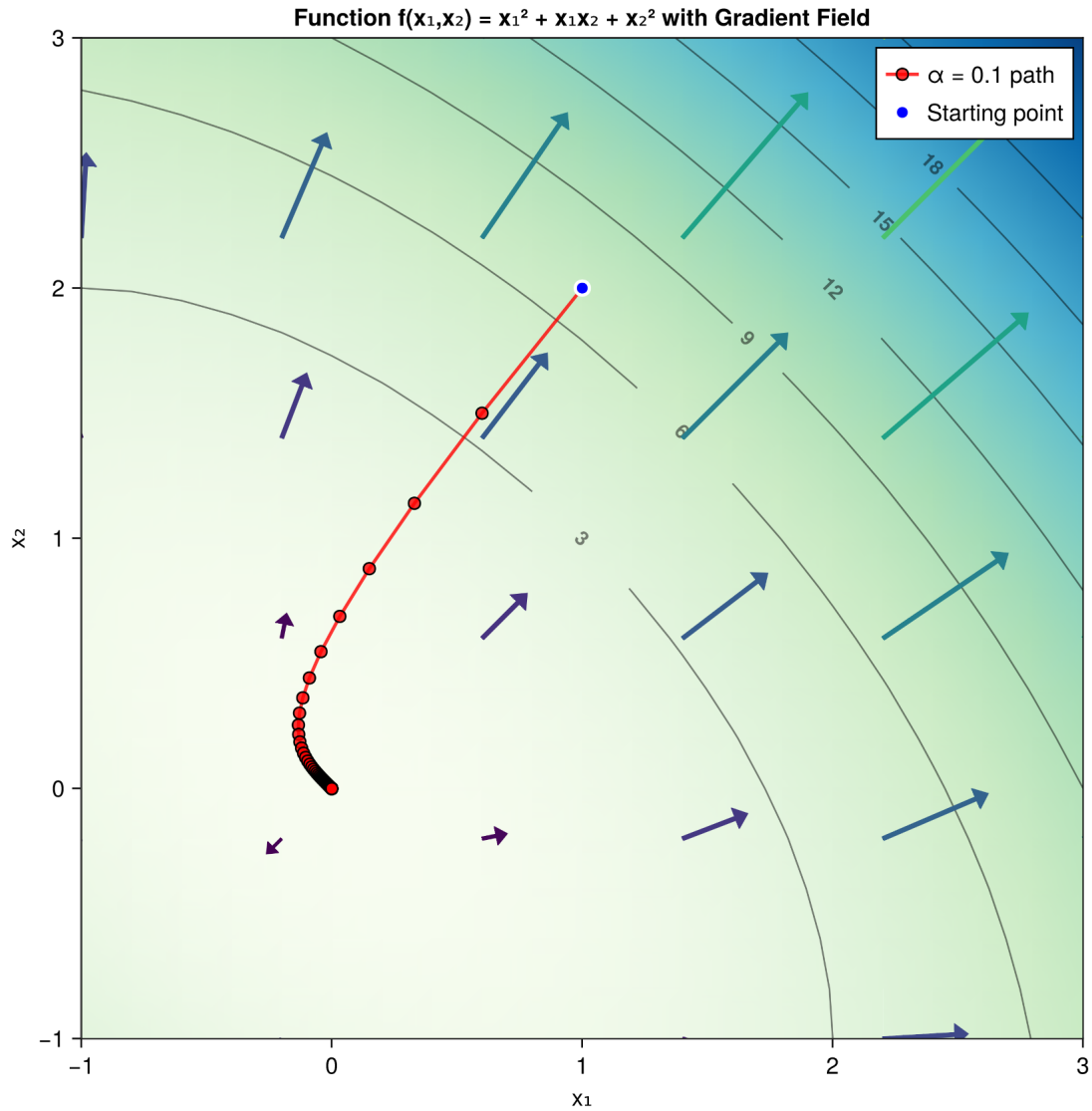


图 3 Gradient Descent Visualization

It is not difficult to see that the iteration speed is related to the step size, which can lead to divergence. Therefore, we need to choose an appropriate step size to ensure convergence.

1.3 CONSTRAINED CASE

When dealing with constrained optimization problems, we need to find the minimum of a function $f(\mathbf{x})$ subject to constraints. The general form is:

$$\begin{aligned} \min_{\mathbf{x} \in \mathbf{R}^n} \quad & f(\mathbf{x}) \\ \text{subject to} \quad & g_i(\mathbf{x}) \leq 0, \quad i = 1, 2, \dots, m \\ & h_j(\mathbf{x}) = 0, \quad j = 1, 2, \dots, l \end{aligned}$$

For constrained problems, we cannot simply move in the negative gradient direction as this may violate the constraints. Instead, we need to project the gradient onto the feasible region or use penalty methods.

1.3.1 PROJECTED GRADIENT METHOD

The projected gradient method modifies the standard gradient descent by projecting each iteration onto the feasible set $\mathcal{C}$:

$$\mathbf{x}_{k+1} = \Pi_{\mathcal{C}}(\mathbf{x}_k - \alpha \nabla f(\mathbf{x}_k))$$

where $\Pi_{\mathcal{C}}$ denotes the projection operator onto the constraint set $\mathcal{C}$.

The pseudocode for the projected gradient method is:

Algorithm Projected Gradient Method for minimize f subject to $\mathbf{x} \in \mathcal{C}$

Input: initial point $\mathbf{x}_0 \in \mathcal{C}$, step length $\alpha > 0$, tolerance $\varepsilon > 0$, max iterations N

Output: $\mathbf{x}^*$

```

1   $k \leftarrow 0$ 
2  while  $k < N$ 
3       $\mathbf{g}_k \leftarrow \nabla f(\mathbf{x}_k)$ 
4       $\mathbf{y}_{k+1} \leftarrow \mathbf{x}_k - \alpha \mathbf{g}_k$ 
5       $\mathbf{x}_{k+1} \leftarrow \Pi_{\mathcal{C}}(\mathbf{y}_{k+1})$ 
6      if  $f(\mathbf{x}_{k+1}) < \varepsilon$  then
7           $\perp$  return  $\mathbf{x}_{k+1}$ 
8      end
9       $k \leftarrow k + 1$ 
10 end
11  $\perp$  return  $\mathbf{x}_k$ 

```

1.3.2 EXAMPLE: LINEAR CONSTRAINT

Consider the optimization problem:

$$\min_{(\mathbf{x}_1, \mathbf{x}_2) \in \mathbf{R}^2} f(\mathbf{x}_1, \mathbf{x}_2)$$

$$\text{subject to } \mathbf{x}_2 = 1$$

where $f(\mathbf{x}_1, \mathbf{x}_2) = \mathbf{x}_1^2 + \mathbf{x}_1 \mathbf{x}_2 + \mathbf{x}_2^2$ (same as the unconstrained case).

The feasible set is the line $\mathcal{C} = \{(\mathbf{x}_1, \mathbf{x}_2) : \mathbf{x}_2 = 1\}$. The unconstrained minimum $(0, 0)$ is not feasible, so we expect the constrained optimum to lie on the constraint.

The gradient is the same as in the unconstrained case.

For the linear constraint $\mathbf{x}_2 = 1$, the projection operation onto the line is:

$$\Pi_{\mathcal{C}}(\mathbf{y}) = \begin{bmatrix} \mathbf{y}_1 \\ 1 \end{bmatrix}$$

Algorithm implementation:

$$\begin{bmatrix} x_1 \\ x_2 \end{bmatrix}_{k+1} = \Pi_{\mathcal{C}} \left(\begin{bmatrix} x_1 \\ x_2 \end{bmatrix}_k - \alpha \begin{bmatrix} 2x_1 + x_2 \\ 2x_2 + x_1 \end{bmatrix}_k \right)$$

We demonstrate the algorithm starting from the same initial point as the unconstrained case:

We set the initial point as $\mathbf{x}_0 = (1.0, 2.0)^T$, which lies off the constraint. The algorithm first projects this point onto the constraint $x_2 = 1$, resulting in $(1.0, 1.0)$. After 1000 iterations, it converges to the constrained optimal point with function value 0.75.

The visualization below shows the optimization trajectory and the projection process. The black line represents the constraint $x_2 = 1$.

For all iterations, the visualization shows (with later iterations becoming more transparent):

- **Red arrows:** gradient steps $-\alpha \nabla f(\mathbf{x}_k)$ from current point to unconstrained update
- **Orange dotted lines:** projection steps from the unconstrained update back to the constraint

This clearly demonstrates how the projected gradient method alternates between taking gradient steps and projecting back to the feasible set. The gradient arrows show both the direction and magnitude of the descent step, while the projection steps ensure feasibility. The path successfully reaches the constrained optimal point $(-0.5, 1.0)$.

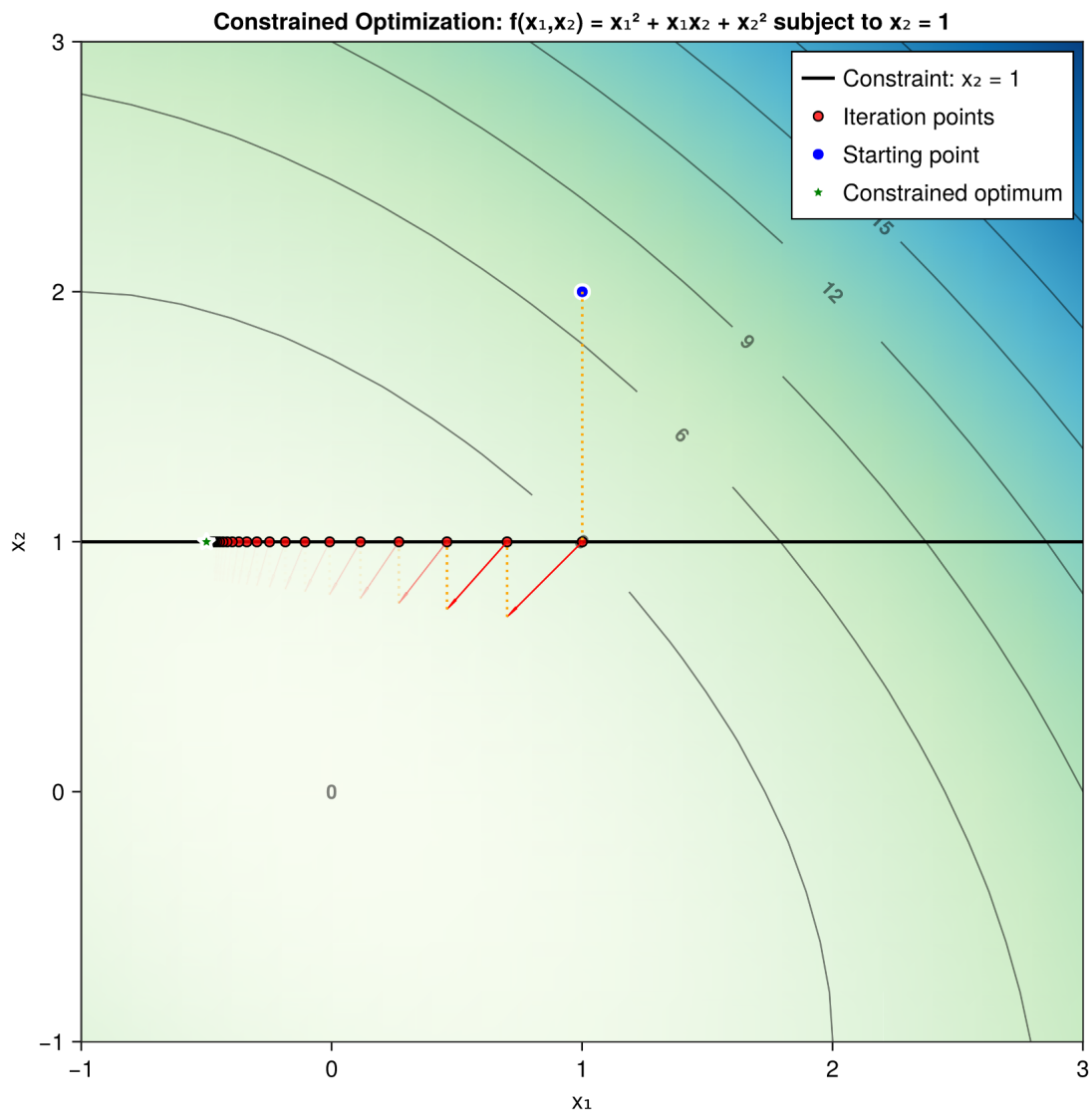


图 4 Constrained Gradient Descent Visualization

1.3.3 PENALTY METHOD

Another approach for handling constraints is the penalty method, where we convert the constrained problem into an unconstrained one by adding penalty terms:

$$\min_{\mathbf{x} \in \mathbf{R}^n} L(\mathbf{x}, \rho) = f(\mathbf{x}) + \rho \sum_{i=1}^m \max(0, g_i(\mathbf{x}))^2 + \rho \sum_{j=1}^l h_{j(\mathbf{x})}^2$$

where $\rho > 0$ is the penalty parameter. As $\rho \rightarrow \infty$, the solution of the penalized problem approaches the solution of the original constrained problem.

1.4 CONVERGENCE

Next, we rigorously analyze the convergence of gradient descent.

1.5 CONVEX OPTIMIZATION

If the optimization problem is convex, then gradient descent can guarantee finding the global minimum. The definition of a convex function is: for any $x, y \in \mathbf{R}^n$ and $\lambda \in [0, 1]$, we have

$$f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y)$$

2 光滑流形上的優化

2.1 幺正流形

幺正群 $U(n)$ 是嵌入在 $M_n(\mathbf{C})$ 上的子流形

2.2 RIEMANNIAN 梯度

Riemannian 梯度是 Euclid 梯度在切空间上的投影. 设 $\bar{f}$ 是 f 到环境空间的光滑延拓, $\text{Proj}_{T_x \mathcal{M}}$ 为到切空间 $T_x \mathcal{M}$ 的正交投影, 则

$$\text{grad } f(x) = \text{Proj}_{T_x \mathcal{M}}(\text{grad } \bar{f}(x)).$$

以幺正流形 $\mathcal{M} = \mathcal{U} = U(n)$ 为例: 切空间是反 Hermitian 方向

$$T_x \mathcal{U} = \{xZ \mid Z^H = -Z\}$$

(Z 取遍反 Hermitian 矩阵, 是集合的哑元), 法空间是 $\{xA : A^H = A\}$, 于是投影给出

$$\text{grad } f(x) = x \text{ skew}(x^H \text{grad } \bar{f}(x)), \quad \text{skew}(A) = \frac{1}{2}(A - A^H).$$

可见 $\text{grad } f(x)$ 形如 xZ , 正是下面回撤所需的切向量输入.

2.3 回撤

DEFINITION 1 回撤映射 流形 $\mathcal{M}$ 上的回撤是光滑映射

$$T\mathcal{M} \ni (x, v) \mapsto R_x(v) \in \mathcal{M}$$

使得任意曲线 $c(t) = R_x(tv)$ 有

$$c(0) = x$$

$$c'(0) = v$$

也就是说, 回撤是指在流形 $\mathcal{M}$ 上的点 x 和切空间 $T_x \mathcal{M}$ 上的点 v 到流形上的映射. 但要保证切空间中自变量沿着 $t \mapsto tv$ 运动时, 回撤到流形上的点 $c(t) = R_x(tv)$ 在起点 $t = 0$ 处沿着 v 方向从 x 出发. 从微分方程的角度来说, 这是在满足初值条件的约束下的曲线. 任意诱导出这样的曲线的映射都是回撤映射.

□ **EXAMPLE 2** 幺正流形上的指数回撤

幺正流形上的点 $x \in \mathcal{M} = \mathcal{U}(n)$. 其切向量 $v \in T_x \mathcal{M}$. $U(n)$ 必形如 $v = x\Omega$, 故

$$\Omega := x^H v \quad (\Omega^H = -\Omega)$$

是由 x, v 确定的常量反 Hermitian 矩阵 (Riemannian 梯度 $\text{grad } f(x) = x \text{ skew}(x^H \text{grad } \bar{f}(x))$ 即是这样一个切向量的例子). 目标: 构造以参数 $t \in \mathbb{R}$ 为唯一自变量、过 x 且初速为 v 的回撤曲线 $c(t)$.

回撤只要求 $c(0) = x$ 、 $c'(0) = v = c(0)\Omega$ 。自然的想法是不只在起点、而是处处令速度满足同一关系 $c'(t) = c(t)\Omega$ 。这便得到一个简单的常系数线性初值问题

$$\begin{cases} c'(t) = c(t)\Omega \\ c(0) = x \end{cases}$$

其唯一解为矩阵指数

$$c(t) = x \exp(t\Omega).$$

还需验证曲线始终留在流形上。令 $P(t) = c(t)^H c(t)$ ，则

$$P'(t) = c'(t)^H c(t) + c(t)^H c'(t) = \Omega^H P + P \Omega,$$

在 $P = I$ 处 $\Omega^H + \Omega = 0$ ，故 $P \equiv I$ ，即 $c(t)$ 全程么正。

因此回撤为

$$R_x(x\Omega) = c(1) = x \exp(\Omega),$$

以切向量 v 为自变量即

$$R_x(v) = x \exp(x^H v).$$

□

上面这条曲线恰好也是 $\mathcal{U}$ 上的测地线。但把它当作回撤**并不需要**测地线性质：我们只用到“满足一阶初值条件”与“留在流形上”两点，测地线只是附带结论。