

Photogrammetry & Robotics Lab

Some Math Basics

Cyrill Stachniss

The slides have been created by Cyrill Stachniss.

Motivation

- We use several concepts from math
- **Goal:** Provide a short reminder for few things that we will use on our way

Brief, informal, incomplete, and unordered set of explanations

Motivation

- We use several concepts from math
- **Goal:** Provide a short reminder
- **Topics**
 - Solving $Ax=b$
 - Solving $Ax=0$ using SVD
 - Least squares with Gauss Newton
 - Skew-symmetric matrix
 - Derivative of rotation matrices
 - Homogenous coordinates (own lecture)

System of Linear Equations

$$Ax = b$$

Linear Equation System: $Ax=b$

Three cases:

- A is squared and has full rank
- A is **over**determined
- A is **under**determined

Solving $Ax=b$, w/ Exact Solution

- **A is a square matrix with full rank**
- Best-case situation, unique solution
- Can be solved in many ways...

Solving $Ax=b$, w/ Exact Solution

- **A is a square matrix with full rank**
- Best-case situation, unique solution
- Can be solved through
 - Gauss elimination
 - Inversion of A : $x = A^{-1}b$
 - Cholesky decomposition $\text{chol}(A) = LL^T$
with lower triangular matrix L
and then solving $Ly = b$ and $L^T x = y$
 - QR decomposition
 - Conjugate gradients

Solving $Ax=b$, A overdetermined

- **Common real-world situation**
- No exact solution exists
- We aim at finding minimizing $\|Ax - b\|$ instead of solving $Ax = b$:

$$x^* = \arg \min_x \|Ax - b\|$$

- Ordinary least squares approach
- Solution can be obtained through

$$x = (A^T A)^{-1} A^T b$$

Solving $Ax=b$, A underdetermined

- **Infinitely many solutions exist**
(or no solution if inconsistent)
- Not enough information available
- Approach: Find x which solves $Ax = b$
and minimizes $\|x\|$
- Solution

$$x = A^T (AA^T)^{-1} b$$

Homogenous System

$$Ax = 0$$

Homogenous System: $Ax=0$

- Find a solution $x \neq 0$ fulfilling $Ax = 0$
- Means system is underdetermined
- There exists a null space of A called $\text{null}(A)$ and all x fulfilling $Ax = 0$ are elements of it
- A 's rank deficiency defines the dimensionality of the null space

Eigenvalues

- For a squared matrix, we have

$$\dim(A) = \dim(\text{null}(A)) + \text{rank}(A)$$

- Which impact does this have on the Eigenvalues of A ?

Eigenvalues

- For a squared matrix, we have

$$\dim(A) = \dim(\text{null}(A)) + \text{rank}(A)$$

- Which impact does this have on the Eigenvalues of A ?
- There are $\text{rank}(A)$ non-zero Eigenvalues
- There are $\dim(\text{null}(A))$ Eigenvalues that are zero

Eigenvector

- For each Eigenvector ν holds $A\nu = \lambda\nu$
- Thus, for those with Eigenvalue 0 we have $A\nu = 0\nu = 0$

Eigenvector

- For each Eigenvector ν holds $A\nu = \lambda\nu$
- Thus, for those with Eigenvalue 0 we have $A\nu = 0\nu = 0$
- **Result:** all Eigenvectors corresponding to an Eigenvalue of 0 solve $Ax = 0$
- The same holds for all linear combinations of these Eigenvectors
- These Eigenvectors form $\text{null}(A)$

Eigenvector & Singular Vectors

- If A is square, real, symmetric and has non-negative Eigenvalues, then Eigenvalues equal to singular values
- Singular vectors and values also defined for non-square matrices
- We can use SVD to compute the singular values and vectors

Singular Value Decomposition

- SVD decomposes a matrix A into

$$A = UDV^T$$

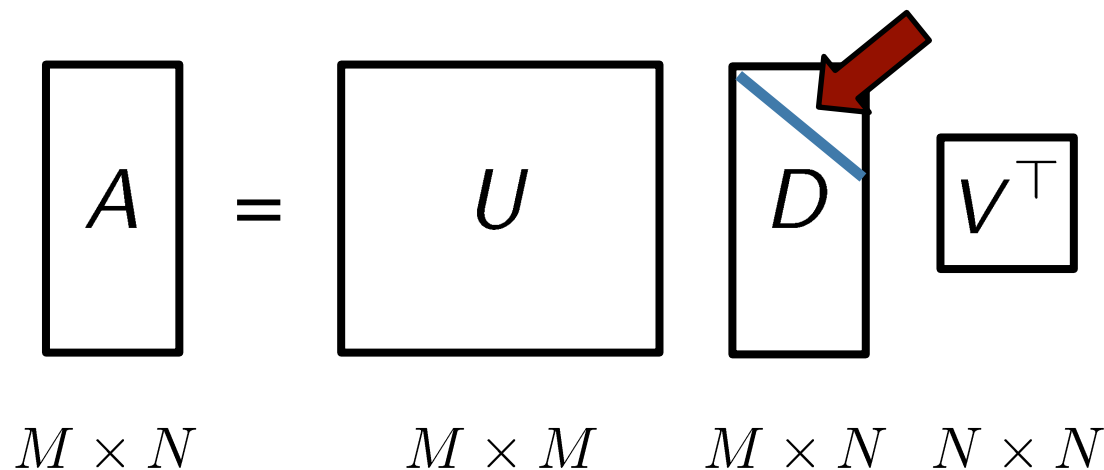
A diagram illustrating the dimensions of the matrices in the SVD decomposition. It shows the equation $A = UDV^T$ where each matrix is enclosed in a box. Below each box is its dimension: A is $M \times N$ (represented by a tall rectangle), U is $M \times M$ (represented by a square), D is $M \times N$ (represented by a tall rectangle), and V^T is $N \times N$ (represented by a small square).

$$\begin{array}{ccccccc} \boxed{A} & = & \boxed{U} & & \boxed{D} & & \boxed{V^T} \\ M \times N & & M \times M & & M \times N & & N \times N \end{array}$$

Singular Values

- SVD decomposes a matrix A into

$$A = UDV^T$$



- D is a diagonal matrix of singular values sorted from large to small
- U, V are orthogonal matrices

Singular Vectors

- SVD decomposes a matrix A into

$$A = UDV^T$$

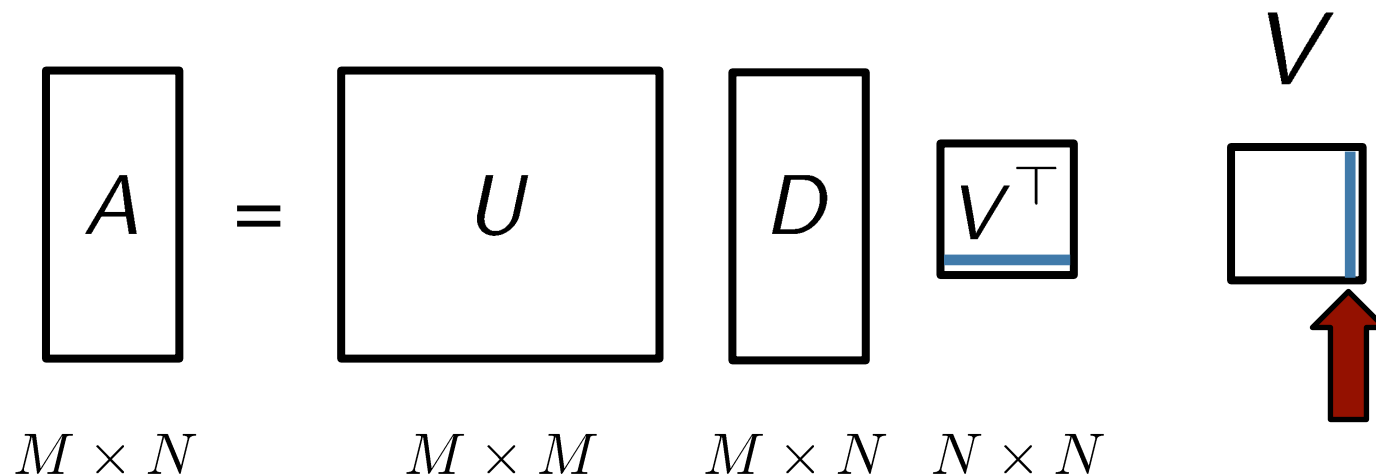
The diagram shows the SVD decomposition $A = UDV^T$ using boxes to represent matrices. Matrix A is a tall rectangle labeled $M \times N$. Matrix U is a square labeled $M \times M$. Matrix D is a tall rectangle labeled $M \times N$. Matrix V^T is a small square labeled $N \times N$. A red arrow points to the bottom row of the V^T box, which is highlighted with a blue line. Below each matrix box is its dimension: $M \times N$ for A , $M \times M$ for U , $M \times N$ for D , and $N \times N$ for V^T .

- V^T stores the corresponding singular vectors to the values

Singular Vectors

- SVD decomposes a matrix A into

$$A = UDV^T$$



- Math libraries often returns V not V^T
- The last column of V stores the vector corresponding to the smallest value

Solution to $Ax=0$ via SVD

- Decompose A using SVD: $A = UDV^T$
- Check if the smallest singular value in D is zero: $D_{NN} \stackrel{?}{=} 0$
- **If so, the last column of V is a non-trivial solution x to $Ax = 0$**

Solution to $Ax=0$ via SVD

- Decompose A using SVD: $A = UDV^T$
- Check if the smallest singular value in D is zero: $D_{NN} \stackrel{?}{=} 0$
- If so, the last column of V is a non-trivial solution x to $Ax = 0$
- **If not**, there is **no non-trivial solution** (i.e., only the trivial exists)
- **However**, the last column of V represents the vector that minimizes $\|Ax\|$ under the constraint $\|x\| = 1$

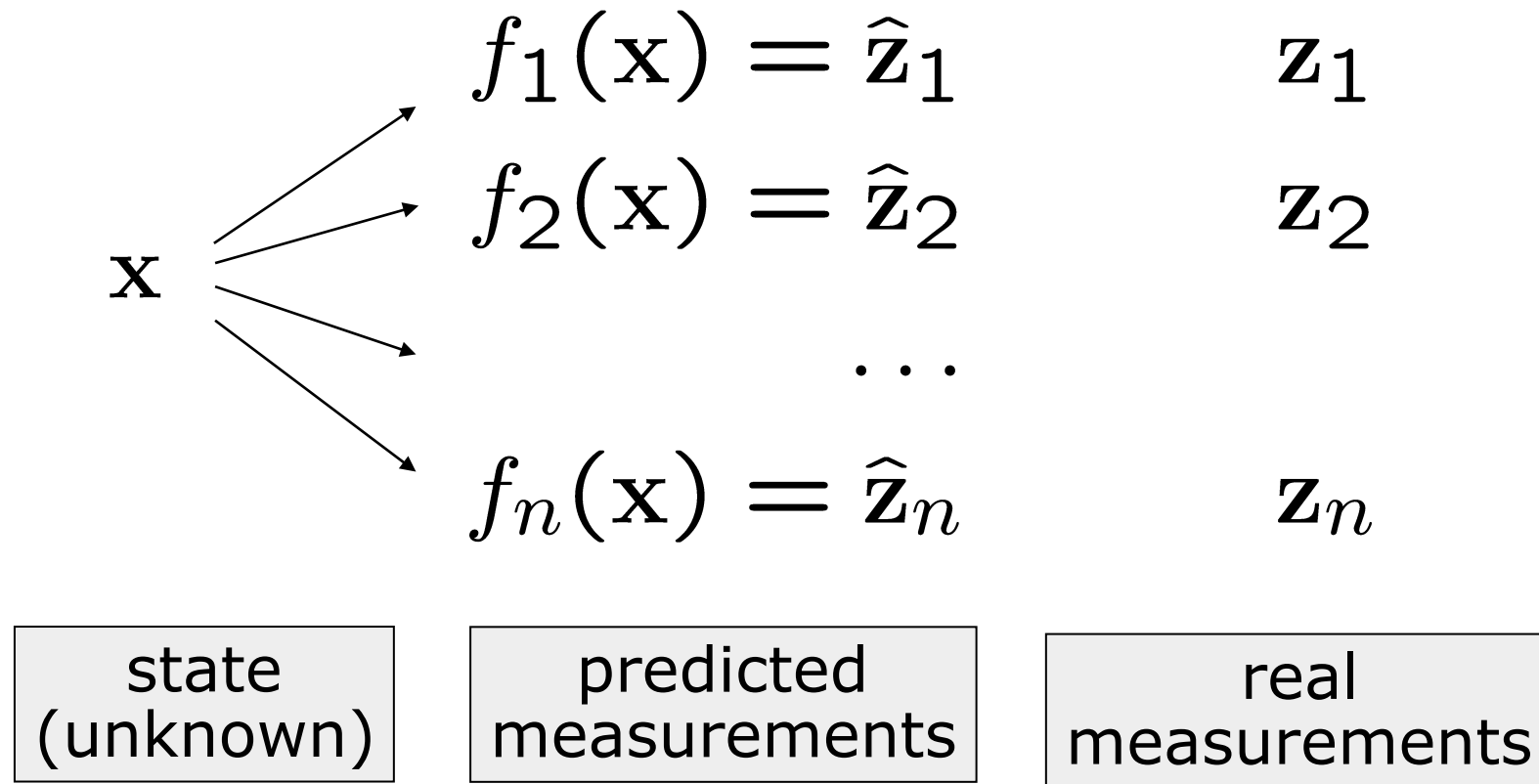
Least Squares (an non-Geodetic view)

Least Squares in 5 Minutes



<https://www.youtube.com/watch?v=87S82fh4rI4>

Graphical Explanation



Error Function

- Error $\mathbf{e}_i$ is typically the **difference** between the **predicted and actual** measurement

$$\mathbf{e}_i(\mathbf{x}) = \mathbf{z}_i - f_i(\mathbf{x})$$

- We assume that the error has **zero mean** and is **normally distributed**
- Gaussian error with information matrix Λ_i
- The squared error of a measurement depends only on the state and is a scalar

$$e_i(\mathbf{x}) = \mathbf{e}_i(\mathbf{x})^T \Lambda_i \mathbf{e}_i(\mathbf{x})$$

Linearizing the Error Function

- Approximate the error functions around an initial guess $\mathbf{x}$ via Taylor expansion

$$e_i(\mathbf{x} + \Delta\mathbf{x}) \simeq \underbrace{e_i(\mathbf{x})}_{e_i} + \mathbf{J}_i(\mathbf{x}) \Delta\mathbf{x}$$

- $\mathbf{J}$ is the Jacobian

$$\mathbf{J}_f(\mathbf{x}) = \begin{pmatrix} \frac{\partial f_1(\mathbf{x})}{\partial x_1} & \frac{\partial f_1(\mathbf{x})}{\partial x_2} & \cdots & \frac{\partial f_1(\mathbf{x})}{\partial x_n} \\ \frac{\partial f_2(\mathbf{x})}{\partial x_1} & \frac{\partial f_2(\mathbf{x})}{\partial x_2} & \cdots & \frac{\partial f_2(\mathbf{x})}{\partial x_n} \\ \cdots & \cdots & \cdots & \cdots \\ \frac{\partial f_m(\mathbf{x})}{\partial x_1} & \frac{\partial f_m(\mathbf{x})}{\partial x_2} & \cdots & \frac{\partial f_m(\mathbf{x})}{\partial x_n} \end{pmatrix}$$

Gauss-Newton

Iterate the following steps:

- Linearize around $\mathbf{x}$ and compute for each measurement

$$e_i(\mathbf{x} + \Delta\mathbf{x}) \simeq e_i(\mathbf{x}) + \mathbf{J}_i \Delta\mathbf{x}$$

- Compute the terms for the linear system $\mathbf{b}^\top = \sum_i e_i^\top \Lambda_i \mathbf{J}_i$ $\mathbf{H} = \sum_i \mathbf{J}_i^\top \Lambda_i \mathbf{J}_i$

- Solve the linear system

$$\Delta\mathbf{x}^* = -\mathbf{H}^{-1}\mathbf{b}$$

- Updating state $\mathbf{x} \leftarrow \mathbf{x} + \Delta\mathbf{x}^*$

Skew-Symmetric Matrices

Skew-Symmetric Matrices

- A skew-symmetric matrix is a matrix S for which holds $S^T = -S$

Skew-Symmetric Matrices

- A skew-symmetric matrix is a matrix S for which holds $S^T = -S$
- S has zeros on the main diagonal
- $\forall S \in \mathbb{R}^{3 \times 3} : \det(S) = 0$
- $\det(S) = 0$ if $\dim(S)$ odd.

Skew-Symmetric Matrices in 3D

- In $\mathbb{R}^3$ we can express the cross product through a skew-symmetric matrix

$$\mathbf{a} \times \mathbf{b} = [\mathbf{a}]_{\times} \mathbf{b} = S_{\mathbf{a}} \mathbf{b}$$

$$[\mathbf{a}]_{\times} = S_{\mathbf{a}} = \begin{bmatrix} 0 & -a_3 & a_2 \\ a_3 & 0 & -a_1 \\ -a_2 & a_1 & 0 \end{bmatrix}$$

Skew-Symmetric Matrices in 3D

- In $\mathbb{R}^3$ we can express the cross product through a skew-symmetric matrix

$$\mathbf{a} \times \mathbf{b} = [\mathbf{a}]_{\times} \mathbf{b} = S_a \mathbf{b}$$

$$[\mathbf{a}]_{\times} = S_a = \begin{bmatrix} 0 & -a_3 & a_2 \\ a_3 & 0 & -a_1 \\ -a_2 & a_1 & 0 \end{bmatrix}$$

$$\underbrace{\begin{bmatrix} a_1 \\ a_2 \\ a_3 \end{bmatrix}}_{\mathbf{a}} \times \underbrace{\begin{bmatrix} b_1 \\ b_2 \\ b_3 \end{bmatrix}}_{\mathbf{b}} = \begin{bmatrix} -a_3b_2 + a_2b_3 \\ a_3b_1 - a_1b_3 \\ -a_2b_1 + a_1b_2 \end{bmatrix} = \underbrace{\begin{bmatrix} 0 & -a_3 & a_2 \\ a_3 & 0 & -a_1 \\ -a_2 & a_1 & 0 \end{bmatrix}}_{S_a} \underbrace{\begin{bmatrix} b_1 \\ b_2 \\ b_3 \end{bmatrix}}_{\mathbf{b}}$$

Derivative of a Rotation Matrix

Derivative of a Rotation Matrix

- Skew-symmetric matrices are useful to formulate the derivative of a rotation matrix

- For any rotation matrix R holds

$$RR^{\top} = I$$

- Consider a rotation by θ around x-axis

$$R_x(\theta)$$

- Then, we have $R_x(\theta)R_x^{\top}(\theta) = I$

Derivative of a Rotation Matrix

- Compute derivative (chain rule)

$$R_x(\theta)R_x^\top(\theta) = I$$

$$\frac{d}{d\theta} \left(R_x(\theta)R_x^\top(\theta) \right) = \frac{d}{d\theta} I$$

$$\frac{d}{d\theta} R_x(\theta)R_x^\top(\theta) + R_x(\theta)\frac{d}{d\theta} R_x^\top(\theta) = 0$$

Derivative of a Rotation Matrix

- Compute derivative (chain rule)

$$R_x(\theta)R_x^\top(\theta) = I$$

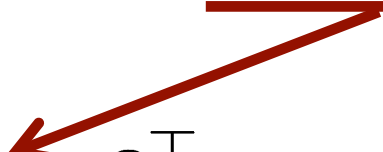
$$\frac{d}{d\theta} \left(R_x(\theta)R_x^\top(\theta) \right) = \frac{d}{d\theta} I$$

$$\frac{d}{d\theta} R_x(\theta)R_x^\top(\theta) + R_x(\theta)\frac{d}{d\theta} R_x^\top(\theta) = 0$$

- Exploiting $(AB)^\top = B^\top A^\top$ leads us to

$$\frac{d}{d\theta} R_x(\theta)R_x^\top(\theta) + \left(\frac{d}{d\theta} R_x(\theta)R_x^\top(\theta) \right)^\top = 0$$

Derivative of a Rotation Matrix

- Rewrite $\frac{d}{d\theta} R_x(\theta) R_x^\top(\theta) + \left(\frac{d}{d\theta} R_x(\theta) R_x^\top(\theta) \right)^\top = 0$ 
- as $S + S^\top = 0$
- This directly leads to $S^\top = -S$, which is a skew-symmetric matrix
- We can now exploit the fact that $\frac{d}{d\theta} R_x(\theta) R_x^\top(\theta)$ is a skew-symmetric matrix

Derivative of a Rotation Matrix

- We have $S = \frac{d}{d\theta} R_x(\theta) R_x^\top(\theta)$

$$R_x = \begin{bmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & \sin \theta \\ 0 & -\sin \theta & \cos \theta \end{bmatrix}$$

- So

$$S = \underbrace{\begin{bmatrix} 0 & 0 & 0 \\ 0 & -\sin \theta & -\cos \theta \\ 0 & \cos \theta & -\sin \theta \end{bmatrix}}_{\frac{d}{d\theta} R_x(\theta)} \underbrace{\begin{bmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & -\sin \theta \\ 0 & \sin \theta & \cos \theta \end{bmatrix}}_{R_x^\top(\theta)}$$

Derivative of a Rotation Matrix

$$\begin{aligned} S &= \begin{bmatrix} 0 & 0 & 0 \\ 0 & -\sin \theta & -\cos \theta \\ 0 & \cos \theta & -\sin \theta \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & -\sin \theta \\ 0 & \sin \theta & \cos \theta \end{bmatrix} \\ &= \begin{bmatrix} 0 & 0 & 0 \\ 0 & 0 & -1 \\ 0 & 1 & 0 \end{bmatrix} \\ &= S e_x \end{aligned}$$

with the unit vector $e_x = \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}$

Derivative of a Rotation Matrix

- This means $\frac{d}{d\theta} R_x(\theta) R_x^\top(\theta) = S e_x$
- and thus

$$\frac{d}{d\theta} R_x(\theta) = \frac{d}{d\theta} R_x(\theta) \underbrace{R_x^\top(\theta) R_x(\theta)}_I = S e_x R_x(\theta)$$

- The derivative of a rotation matrix $R_x(\theta)$ is the skew-symmetric matrix $S e_x$ times the rotation matrix itself

$$\frac{d}{d\theta} R_x(\theta) = S e_x R_x(\theta)$$

The Same for x,y,z Axes

- We can repeat the same to x, y, z and obtain

$$\frac{d}{d\theta} R_x(\theta) = S_{e_x} R_x(\theta)$$

$$\frac{d}{d\theta} R_y(\theta) = S_{e_y} R_y(\theta)$$

$$\frac{d}{d\theta} R_z(\theta) = S_{e_z} R_z(\theta)$$

- and even for an arbitrary rot. axis r

$$\frac{d}{d\theta} R_r(\theta) = S_r R_r(\theta)$$

Infinitesimal Small Rotations

- Similarly, we can also approximate an infinitesimally small rotation by

$$R \approx I + dR = I + S_{dr} = I + \begin{bmatrix} 0 & -d\kappa & d\phi \\ d\kappa & 0 & -d\omega \\ -d\phi & d\omega & 0 \end{bmatrix}$$

- Thus,

$$dR = S_{dr} = \begin{bmatrix} 0 & -d\kappa & d\phi \\ d\kappa & 0 & -d\omega \\ -d\phi & d\omega & 0 \end{bmatrix}$$

Summary

This lecture was a **brief and informal reminder** of concepts we will need

- Solving $Ax=b$
- Solving $Ax=0$ using SVD
- Least squares with Gauss Newton
- Skew-symmetric matrices
- Derivative of a rotation matrix
- Own lecture: Homogenous coordinates

Slide Information

- The slides have been created by Cyrill Stachniss as part of the photogrammetry and robotics courses.
- **I tried to acknowledge all people from whom I used images or videos. In case I made a mistake or missed someone, please let me know.**
- The photogrammetry material heavily relies on the very well written lecture notes by Wolfgang Förstner and the Photogrammetric Computer Vision book by Förstner & Wrobel.
- Parts of the robotics material stems from the great Probabilistic Robotics book by Thrun, Burgard and Fox.
- If you are a university lecturer, feel free to use the course material. If you adapt the course material, please make sure that you keep the acknowledgements to others and please acknowledge me as well. To satisfy my own curiosity, please send me email notice if you use my slides.