

量化投资：以 Python 为工具

蔡立岗

第一部分

基础

第一章 略。

第二章 Python 代码的编写与执行

1. ¹

```
1 In [1]: %paste
2 for i in range(1,11):
3     if i%2 == 1:
4         print(i)
5
6 ## — End pasted text —
7 1
8 3
9 5
10 7
11 9
```

¹在代码文件中复制代码，然后在 `ipython` 中输入 `%paste`，并按 `enter` 键，即可执行代码并输出结果。

第三章 Python 对象类型初探

1.

```
1 In [1]: type((1,3,5,7))
2 Out[1]: tuple
3
4 In [2]: type('p')
5 Out[2]: str
6
7 In [3]: type(5>6)
8 Out[3]: bool
9
10 In [4]: type({'High':7.45,'Low':7.30})
11 Out[4]: dict
12
13 In [5]: type(['Hello world',3,4,('a','b'),True])
14 Out[5]: list
```

2.

```
1 In [1]: (3+4j)*(4+3j)
2 Out[1]: 25j
3
4 In [2]: 3/4
5 Out[2]: 0.75
6
7 In [3]: True *3
8 Out[3]: 3
9
10 In [4]: 0.003-0.0022222
11 Out[4]: 0.0007777999999999999
```

3.

```
1 In [1]: a = [1, 2, 3]
2
```

```

3 In [2]: b=a
4
5 In [3]: b
6 Out[3]: [1, 2, 3]
7
8 In [4]: id(a)==id(b)
9 Out[4]: True
10
11 In [5]: b[0]=3
12
13 In [6]: a
14 Out[6]: [3, 2, 3]
15
16 In [7]: id(a)==id(b)
17 Out[7]: True

```

4.

(a)

```

1 In [1]: a=tuple(range(1,11))

```

```

1 In [2]: b=tuple((i+1 for i in range(20) if (i+1)%2==1))

```

(b)

```

1 In [3]: c=[5*i for i in range(11)]

```

```

1 In [4]: %paste
2 d=[];
3 for i in range(1,6):
4     j=0
5     while(j<3):
6         d.append(i)
7         j+=1
8 d
9
10 ## — End pasted text —
11 Out[4]: [1, 1, 1, 2, 2, 2, 3, 3, 3, 4, 4, 4, 5, 5, 5]

```

(d)

```

1 In [5]: e=set({' NASDAQ ', 'Dowjones', ' DAX ', ' FTSE ' })

```

```

1 In [1]: tuple(' abc ')
2 Out[1]: (' a ', ' b ', ' c ')

```

```
3
4 In [2]: list(' abc ')
5 Out[2]: [' a ', ' b ', ' c ']
6
7 In [3]: set(' abc ')
8 Out[3]: {' a ', ' b ', ' c '}
```

```
5.
1 In [1]: characters={ A :ord(' A '), b :ord(' b '), \n :ord(' \n ')}
2 In [2]: characters
3 Out[2]: { A : 65, ' \n ' : 10, ' b ' : 98}
```

第四章 略。

第五章 Python 运算符与使用

1.

```
1 In [1]: 3+4
2 Out[1]: 7
3
4 In [2]: 3.4*4
5 Out[2]: 13.6
6
7 In [3]: 3//4
8 Out[3]: 0
9
10 In [4]: id(['d',3,True,'scd'][1])== id(3)
11 Out[4]: True
12
13 In [5]: 3==4 in [1,'345',3+4j,4 in [1,2,3]]
14 Out[5]: False
15
16 In [6]: (3==4*4.5%2 is 0) in [3,4,'Tom','d in 'comic']
17 Out[6]: False
```

2. is 比较的是两个变量所指向的对象是否一样，== 判断的是值是否一样。

3.

```
1 In [1]: a=[1,2,3]
2 In [2]: c=2
3 In [3]: sum([i>c for i in a])==len(a)
4 Out[3]: False
```

```
1 In [1]: import random
2 #因产生随机数，代码结果可能有所不同
3 In [2]: a=[random.normalvariate(0,1) for i in range(20)]
4
5 In [3]: a.sort()
6
7 In [4]: a[0]
```

```
8 Out[4]: -1.5288923718112197
9
10 In [5]: a[-1]
11 Out[5]: 2.2531121988580143
12
13 In [6]: sum(a)
14 Out[6]: 5.7283863023842905
```

第六章 Python 常用语句

4.

```
1 In [1]: %paste
2 def is_odd(i):
3     if type(i) == int:
4         return(i if i%2==1 else 0)
5     else:
6         print( 'ERROR: Please input an integer.' )
7 is_odd(7)
8
9 ## — End pasted text —
10 Out[1]: 7
```

```
1 In [2]: %paste
2 import random
3 a=[random.normalvariate(0,1) for i in range(5)]
4 for j in a:
5     if j>=0:
6         print( Big )
7     else:
8         print( Small )
9
10 ## — End pasted text —
11 Small
12 Small
13 Big
14 Small
15 Big
```

2.

```
1 In [3]: %paste
2 l = []
3 for i in range(5):
4     row = [0 for j in range(5)]
5     row[i] = 1
```

```
6     l.append(row)
7 l
8
9 ## — End pasted text —
10 Out[5]:
11 [[1, 0, 0, 0, 0],
12  [0, 1, 0, 0, 0],
13  [0, 0, 1, 0, 0],
14  [0, 0, 0, 1, 0],
15  [0, 0, 0, 0, 1]]
```

```
1 In [4]: for i in (-1,0,1,2,39):
2     ...:     print(i in range(1,5))
3     ...:
4     ...:
5 False
6 False
7 True
8 True
9 False
```

```
4.
1 In [5]: %paste
2 Hilbert = []
3 for i in range(4):
4     Hilbert.append([])
5     for j in range(4):
6         Hilbert[i].append(1/(i+j+1))
7 Hilbert
8
9 ## — End pasted text —
10 Out[6]:
11 [[1.0, 0.5, 0.3333333333333333, 0.25],
12  [0.5, 0.3333333333333333, 0.25, 0.2],
13  [0.3333333333333333, 0.25, 0.2, 0.16666666666666666],
14  [0.25, 0.2, 0.16666666666666666, 0.14285714285714285]]
```

第七章 函数

```
1.
1 In [1]: def permutation(x):
2     ....:     x[0],x[-1]=x[-1],x[0]
3     ....:     return(x)
4     ....: y=permutation(x)
5     ....: y is x
6     ....:
7 Out[1]: True
```

```
1 In [2]: def sum_lists(x,y):
2     ....:     return([x[i]+y[i] for i in range(len(x))])
3     ....:
4     ....:
```

```
2.
1 In [3]: def sum2(*lists):
2     ....:     if len(lists)==2:
3     ....:         return(sum_lists(lists[0],lists[1]))
4     ....:     else:
5     ....:         return(sum_lists(lists[0],sum2(*lists[1:])))
6     ....:
7     ....:
```

```
1 In [4]: %paste
2 def fibo(n):
3     if(n<3):
4         a=1
5     else:
6         a=fibo(n-1)+fibo(n-2)
7     return(a)
8 def seqfibo(n):
9     for i in range(1,n+1):
10         print(fibo(i))
11 seqfibo(5)
```

```
12
13 ## — End pasted text —
14 1
15 1
16 2
17 3
18 5
```

```
5.
1 In [5]: %paste
2 def multi_abs(*numbers):
3     return(list(map(abs,numbers)))
4 multi_abs(-5,55,-6,0,-7)
5
6 ## — End pasted text —
7 Out[7]: [5, 55, 6, 0, 7]
```

```
1 In [6]: %paste
2 def count_positive(returns):
3     return(sum(map(lambda x: x>0,returns)))
4 ret=[-0.4,0.5,-0.34,0.45,0.50]
5 count_positive(ret)
6
7 ## — End pasted text —
8 Out[9]: 3
```

第八章 面向对象

6.

```
1 In [1]: class account:
2     ...:     def __init__(self,name,balance):
3     ...:         self.name=name
4     ...:         self.balance=balance
5     ...:
6     ...:     def deposit(self,amount):
7     ...:         self.balance += amount
8     ...:
9     ...:     def withdraw(self,amount):
10    ...:         if amount>self.balance:
11    ...:             print(' 余额不足，交易失败 ')
12    ...:         else:
13    ...:             self.balance -= amount
14    ...:
15    ...:
```

(a)

```
1 In [2]: sam = account(' Sam ',1000)
```

```
1 In [3]: sam.deposit(500)
```

```
2 In [4]: sam.withdraw(1200)
```

(b)

```
1 In [5]: sam.balance
```

```
2 Out[5]: 300
```

```
1 In [6]: class account:
2     ...:     def __init__(self,name,balance):
3     ...:         self.name=name
4     ...:         self.__balance=balance
5     ...:
6     ...:     def deposit(self,amount):
7     ...:         self.__balance += amount
```

```
8     ...:
9     ...:     def withdraw(self,amount):
10    ...:         if amount>self.balance:
11    ...:             print( 余额不足, 交易失败 )
12    ...:         else:
13    ...:             self.__balance -= amount
14    ...:
15    ...:     def get_balance(self):
16    ...:         return(self.__balance)
17    ...:
18    ...:     def set_balance(self,amount):
19    ...:         self.__balance = amount
20    ...:
21    ...:
```

```
1 In [7]: class account:
2     ...:     def __init__(self,name,balance):
3     ...:         self.name=name
4     ...:         self.__balance=balance
5     ...:
6     ...:     def deposit(self,amount):
7     ...:         self.__balance += amount
8     ...:
9     ...:     def withdraw(self,amount):
10    ...:         if amount>self.__balance:
11    ...:             print( 余额不足, 交易失败 )
12    ...:         else:
13    ...:             self.__balance -= amount
14    ...:
15    ...:     def get_balance(self):
16    ...:         return(self.__balance)
17    ...:
18    ...:     def set_balance(self,amount):
19    ...:         self.__balance = amount
20    ...:
21    ...:     def transfer(self,amount,target):
22    ...:         if amount>self.__balance:
23    ...:             print( 余额不足, 交易失败 )
24    ...:         else:
25    ...:             target.deposit(amount)
26    ...:             self.withdraw(amount)
```

```

27     ...:
28     ...:
29
30 In [8]: sam = account(' Sam' ,1000)
31
32 In [9]: john = account(' John' ,3000)
33
34 In [10]: john.transfer(1000,sam)

```

(a)

```

1 In [11]: class check(account):
2     ...:     def __init__(self,name,balance,credit):
3     ...:         self.name=name
4     ...:         self.__balance=balance
5     ...:         self.credit = credit
6     ...:         self.overdraft = 0
7     ...:
8     ...:     def deposit(self,amount):
9     ...:         if self.overdraft > 0 and self.overdraft-amount < 0:
10    ...:             self.__balance = amount - self.overdraft
11    ...:             self.overdraft = 0
12    ...:         elif self.overdraft > 0:
13    ...:             self.overdraft -= amount
14    ...:         else:
15    ...:             self.__balance += amount
16    ...:
17    ...:     def withdraw(self,amount):
18    ...:         if self.__balance==0 and amount + self.overdraft > self.credit:
19    ...:             print(' 超出信用额度，交易失败' )
20    ...:             return(1)
21    ...:         elif self.__balance==0:
22    ...:             self.__overdraft += amount
23    ...:             return(0)
24    ...:         elif 0 < self.__balance < amount and amount - self.__balance <=
        self.credit:
25    ...:             self.overdraft = amount - self.__balance
26    ...:             self.__balance=0
27    ...:             return(0)
28    ...:         elif 0 < self.__balance < amount:
29    ...:             print(' 超出信用额度，交易失败' )

```

```
30         ...:         return(1)
31         ...:     else:
32         ...:         self.__balance -= amount
33         ...:         return(0)
34         ...:
35         ...:     def transfer(self,amount,target):
36         ...:         result = self.withdraw(amount)
37         ...:         if result == 0:
38         ...:             target.deposit(amount)
39         ...:
40         ...:
41
42 In [12]: sam = check('Sam',1000,1000)
43
44 In [13]: sam.withdraw(700)
45 Out[13]: 0
46
47 In [14]: sam.withdraw(1500)
48 超出信用额度，交易失败
49 Out[14]: 1
```

第九章 Python 标准库与数据操作

```
1.
1 In [1]: import datetime as dt
2
3 In [2]: now = dt.datetime.now()
4
5 In [3]: now.strftime( '%Y-%m-%d ' )
6 Out[3]: ' 2017-01-17'
```

```
1 In [4]: def create_name():
2     ...:     name = input( ' 输入用户名\n' )
3     ...:     if ord(name[0])<65 or ord(name[0])>122:
4     ...:         print( ' 用户名必须以字母开头' )
5     ...:         return(create_name())
6     ...:     else:
7     ...:         return(name)
8     ...:
9     ...:
10
11 In [5]: def verify_passwd(passwd):
12     ...:     head = 65 <= ord(passwd[0]) <= 122
13     ...:     contain_number = any([str(x) in passwd for x in range(9)])
14     ...:     contain_symbol = any([str(x) in passwd for x in ( ' _ ; * ; # ' )])
15     ...:     return(head and (contain_number or contain_symbol))
16     ...:
17     ...:
18
19 In [6]: def create_passwd():
20     ...:     passwd = input( ' 输入密码\n' )
21     ...:     is_legal = verify_passwd(passwd)
22     ...:     if is_legal:
23     ...:         return(passwd)
24     ...:     else:
25     ...:         return(create_passwd())
26     ...:
```

```

27     ...:
28
29 In [7]: def create_account():
30     ...:     create_name()
31     ...:     create_passwd()
32     ...:     print( 用户创建成功 )
33     ...:
34     ...:
35 In [8]: create_account()
36 输入用户名
37 python
38 输入密码
39 ffw342-rw
40 用户创建成功

```

2.

```

1 In [10]: evens = [i+1 for i in range(100) if (i+1)%2==0]
2 #output省略

```

4. (a)

```

1 In [1]: dates = [dt.datetime(2015,1,13)+dt.timedelta(i) for i in range(5)]
2
3 In [2]: closes = [7.31,7.28,7.40,7.43,7.41]
4
5 In [3]: prices = {dates[i]:closes[i] for i in range(5)}

```

```

1 In [4]: prices[dt.datetime(2015,1,20)] = 7.44
2
3 In [5]: dates.append(dt.datetime(2015,1,20))

```

(b)

```

1 In [16]: prices[dt.datetime(2015,1,21)-dt.timedelta(4)]
2 Out[6]: 7.41

```

```

1 In [7]: prices[dt.datetime(2015,1,16)] = 7.50

```

(d)

```

1 In [8]: %paste
2 import math
3 cash = 10000
4 share = {dates[0]:0}
5 for i in range(1,6):
6     if prices[dates[i]] > prices[dates[i-1]]:
7         buyshare = math.floor(0.5*cash/prices[dates[i]])

```

```
8         share[dates[i]] = buyshare
9         cash = cash - buyshare * prices[dates[i]] + share[dates[i-1]] * prices[dates[i]]
10     else:
11         share[dates[i]] = 0
12 share
13
14 ## — End pasted text —
15 Out[8]:
16 {datetime.datetime(2015, 1, 13, 0, 0): 0,
17  datetime.datetime(2015, 1, 14, 0, 0): 0,
18  datetime.datetime(2015, 1, 15, 0, 0): 675,
19  datetime.datetime(2015, 1, 16, 0, 0): 333,
20  datetime.datetime(2015, 1, 17, 0, 0): 0,
21  datetime.datetime(2015, 1, 20, 0, 0): 508}
```

第十章 常用第三方库：Numpy 库与多维数组

```
1.
1 In [1]: import numpy as np
2
3 In [2]: hilbert = 1 / (np.arange(4)+np.arange(1,5)[: ,np.newaxis])

1 In [1]: %paste
2 import datetime as dt
3 import numpy as np
4 import datetime as dt
5 import math
6 dates = [dt.datetime(2015,1,13)+dt.timedelta(i) for i in range(5)]
7 closes = [7.31,7.28,7.40,7.43,7.41]
8 prices = {dates[i]:closes[i] for i in range(5)}
9 prices[dt.datetime(2015,1,20)] = 7.44
10 dates.append(dt.datetime(2015,1,20))
11 cash = 10000
12 share = {dates[0]:0}
13 for i in range(1,6):
14     if prices[dates[i]]>prices[dates[i-1]]:
15         buyshare = math.floor(0.5*cash/prices[dates[i]])
16         share[dates[i]]= buyshare
17         cash=cash-buyshare*prices[dates[i]]+share[dates[i-1]]*prices[dates[i]]
18     else:
19         share[dates[i]]= 0
20 buyDates=[np.datetime64(date) for date in share.keys() if share[date]>0]
21 buyDates
22
23 ## — End pasted text —
24 Out[1]:
25 [numpy.datetime64(' 2015-01-15T00:00:00.000000 '),
26  numpy.datetime64(' 2015-01-20T00:00:00.000000 '),
27  numpy.datetime64(' 2015-01-16T00:00:00.000000 ')]
28
29 In [2]: %paste
```

```
30 [prices[dt.datetime.  
31    .strptime(np.datetime_as_string(date)[:10],  
32     "%Y-%m-%d' )] for date in buyDates]  
33  
34 ## — End pasted text —  
35 Out[2]: [7.4, 7.44, 7.43]
```

```
2.  
1 In [1]: cosin = np.cos(np.linspace(0,2*np.pi,1001))
```

```
1 In [1]: sample = np.array([0.5,1.43,-1.36,-0.16,0.29,-0.59,  
2     ....:                    1.16,-0.33,0.07,-1.36])  
3  
4 In [2]: np.mean(sample)  
5 Out[2]: -0.0350000000000000031  
6  
7 In [3]: np.var(sample)  
8 Out[3]: 0.789904999999999986  
9  
10 In [4]: np.std(sample)  
11 Out[4]: 0.88876599844953552  
12  
13 In [5]: np.median(sample)  
14 Out[5]: -0.044999999999999998
```

第十一章 常用第三方库：Pandas 与数据处理

4.

```
1 In [1]: import numpy as np
2
3 In [2]: odds = [i+1 for i in range(20) if (i+1)%2==1]
4
5 In [3]: odds = np.array(odds)
```

2.

```
1 In [4]: three = tuple([i+1 for i in range(40) if (i+1)%3==0])
2
3 In [5]: three = np.array(three)
```

```
1 In [6]: same_number = odds[np.in1d(odds,three)]
2
3 In [7]: print(same_number[:round(len(same_number)/2)])
4 [ 3  9]
5
6 In [8]: for i in range(round(len(same_number)/2)):
7     ...:     print(same_number[i])
8     ...:
9     ...:
10 3
11 9
```

3.

```
1 In [9]: np.random.uniform(0,10,10)
2 Out[9]: array([ 9.84971956,  5.63717566,  2.67350085,  3.43081329,  4.73561598,
3           7.8112453 ,  0.92119308,  7.00992592,  6.31048269,  6.949389  ])
```

5.

```
1 In [10]: even = np.arange(2,22,2)
2
3 In [11]: even[-5:]
4 Out[11]: array([12, 14, 16, 18, 20])
5
6 In [12]: even[len(even)-5:]
```

```
7 Out[12]: array([12, 14, 16, 18, 20])
```

```
6.
1 In [1]: %paste
2 import pandas as pd
3 data = pd.DataFrame({' id ':[' a ' ; ' b ' ; ' c ' ; ' d ' ; ' e ' ; ' f ' ],' name ':
4     [' Alice ' ; ' Bob ' ; ' Charlie ' ; ' David ' ; ' Esther ' ; ' Fanny ' ],' age ':
5     [34,36,30,29,32,36]})
6
7 data.T.ix[2]
8
9 ## — End pasted text —
10 Out[1]:
11 0      Alice
12 1        Bob
13 2    Charlie
14 3      David
15 4     Esther
16 5      Fanny
17 Name: name, dtype: object
```

```
7.
1 In [2]: %paste
2 new = pd.DataFrame({' id ':[' g ' ],' name ':[' John ' ],' age ':[19]})
3 data = data.append(new)
4 data.index = data.age
5 data
6
7 ## — End pasted text —
8 Out[2]:
9      age id      name
10 age
11 34    34  a    Alice
12 36    36  b      Bob
13 30    30  c  Charlie
14 29    29  d    David
15 32    32  e    Esther
16 36    36  f    Fanny
17 19    19  g     John
18
19 In [3]: %paste
20 data = data.drop(30)
```

```
21 data
22
23 ## — End pasted text —
24 Out[3]:
25      age id    name
26 age
27 34    34  a    Alice
28 36    36  b      Bob
29 29    29  d    David
30 32    32  e    Esther
31 36    36  f    Fanny
32 19    19  g     John
```

第十二章 常用第三方库：Matplotlib 库与数据可视化

```
1. In [1]: %paste
2
3 import pandas as pd
4 import numpy as np
5 import matplotlib.pyplot as plt
6
7 Money = pd.read_csv(' Data/Part1/012/Money.csv' ,
8                     index_col= 'date' )
9
10 axis1 = plt.subplot()
11 axis1.plot(Money.index,Money.m,'-')
12 plt.title(' Money Supply of Canada ')
13 plt.xlabel(' Year' )
14 axis2 = axis1.twinx()
15 axis2.plot(Money.y,'-')
16
17 ## — End pasted text —
18 Out[1]: [<matplotlib.lines.Line2D at 0x7f223a2e9b70>]
19
20 In [2]: plt.show()
```

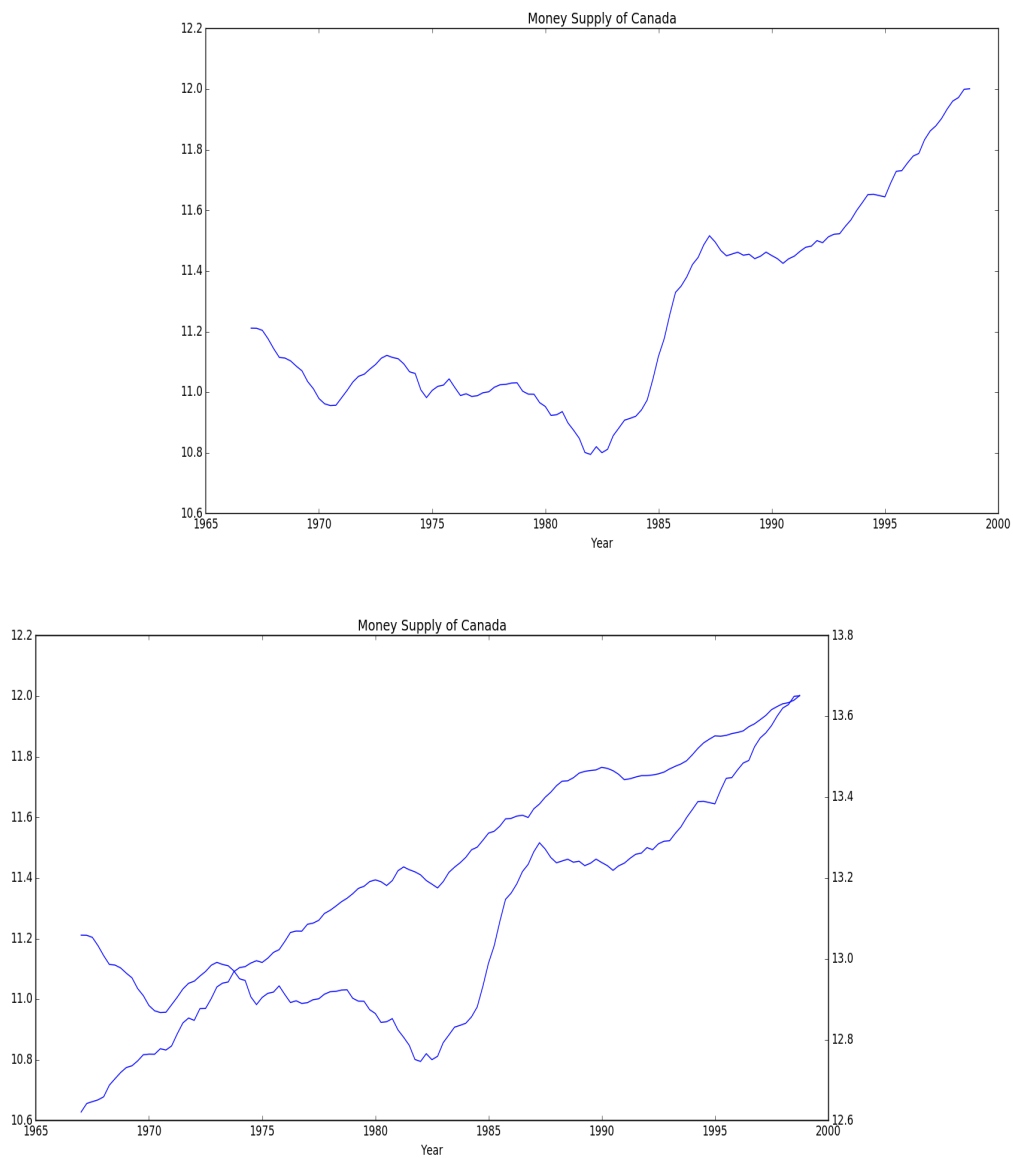


图 12.1: 题 1

```

2. In [1]: %paste
3 Journals = pd.read_csv( Data/Part1/012/Journals.csv )
4 plt.scatter( Journals.citestot, Journals.libprice )
5 plt.title( 'Price vs Citations' )
6 plt.xlabel( 'Citations' )
7 plt.ylabel( 'Price' )
8 plt.show()

```

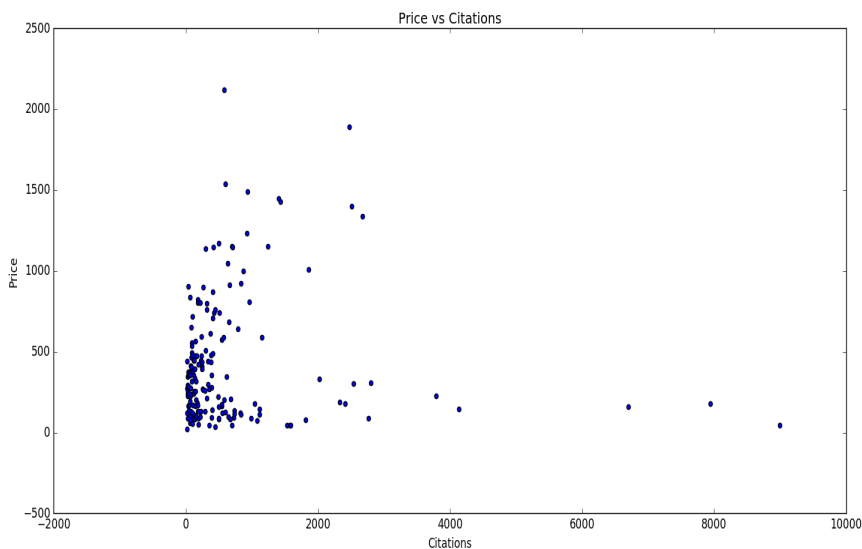


图 12.2: 题 2

```

3.
1 In [1]: %paste
2 mtcars = pd.read_csv(' Data/Part1/012/mtcars.csv ')
3 types=np.array([0.2,0.6])
4 number = mtcars.groupby([' gear' , vs ])[' vs' ].agg(len)
5 plt.bar(types,number.ix[3],width=0.3,label= gear=3 )
6 plt.bar(types,number.ix[4],width=0.3,
7         bottom=number.ix[3],color= ' r' ,
8         label= gear=4 )
9 plt.bar(types,number.ix[5],width=0.3,
10         bottom=number.ix[3]+number.ix[4],
11         color= ' y' ,label= gear=5 )
12 plt.xlim([0,1])
13 plt.ylim([0,25])
14 plt.legend()
15 plt.xticks(types+0.3/2,[0,1])
16 plt.show()

```

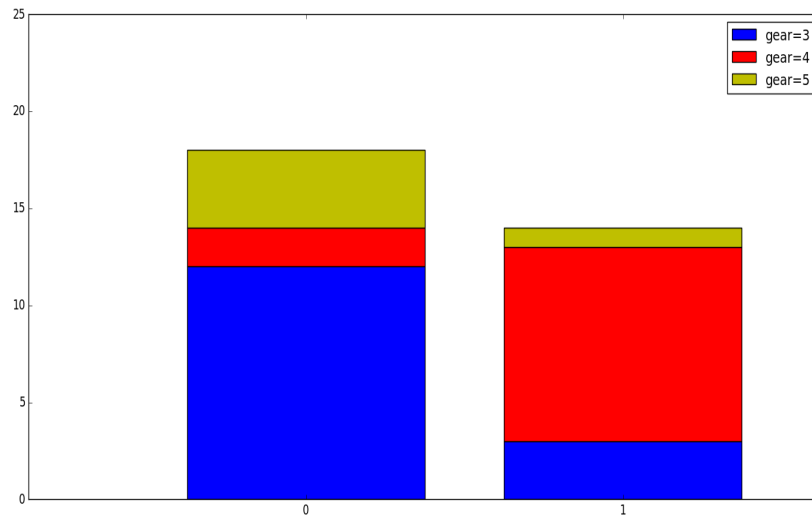


图 12.3: 题 3

```
4. In [1]: %paste
1 Arthritis = pd.read_csv(' Data/Part1/012/Arthritis.csv ')
2 plt.hist(Arthritis.Age)
3 plt.xlabel(' Age ')
4 plt.title(' Histogram of Age ')
5 plt.show()
6 ## — End pasted text —
```

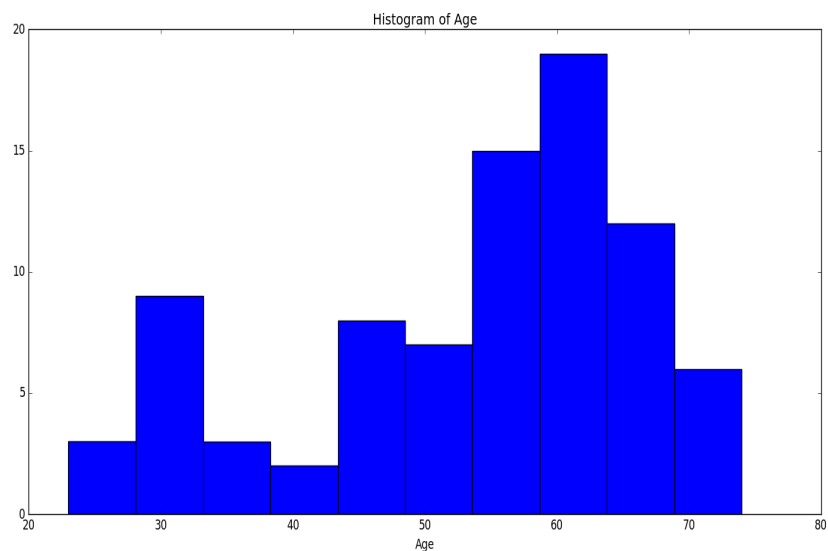


图 12.4: 题 4

```
5. In [1]: %paste
1 norm1 = np.random.normal(0,1,100)
```

```

3 norm2 = np.random.normal(0,2,100)
4 norm3 = np.random.normal(0,3,100)
5 norm4 = np.random.normal(0,4,100)
6 plt.boxplot([norm1,norm2,norm3,norm4])
7 plt.xlabel(' Standard Deviation' )
8 plt.title(' Normal Distributions with Different Standard Deviation' )
9 plt.show()
10 ## — End pasted text —

```

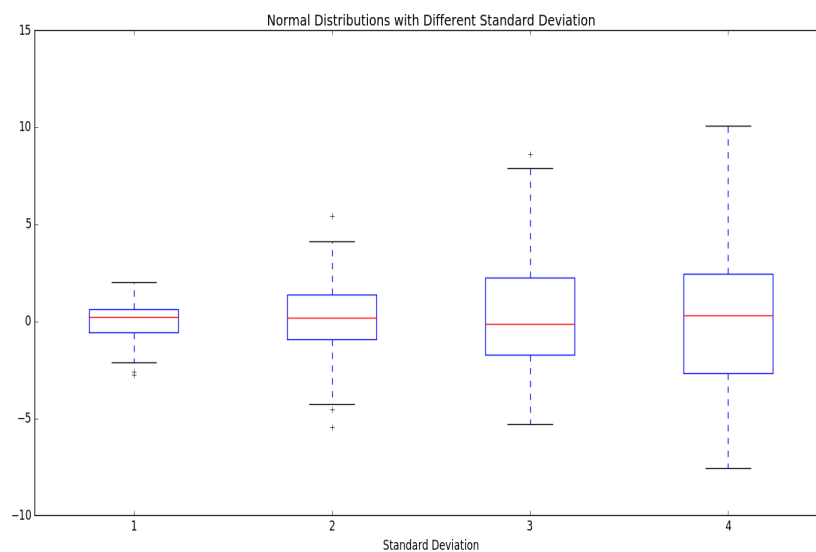


图 12.5: 题 5

```

6. In [1]: %paste
7 norm1 = np.random.normal(0,1,100)
8 norm2 = np.random.normal(0,2,100)
9 norm3 = np.random.normal(0,3,100)
10 norm4 = np.random.normal(0,4,100)
11 figure, axes = plt.subplots(2,2)
12 axes[0,0].scatter(range(100),norm1)
13 axes[0,1].scatter(range(100),norm2)
14 axes[1,0].scatter(range(100),norm3)
15 axes[1,1].scatter(range(100),norm4)
16 plt.show()
17 ## — End pasted text —

```

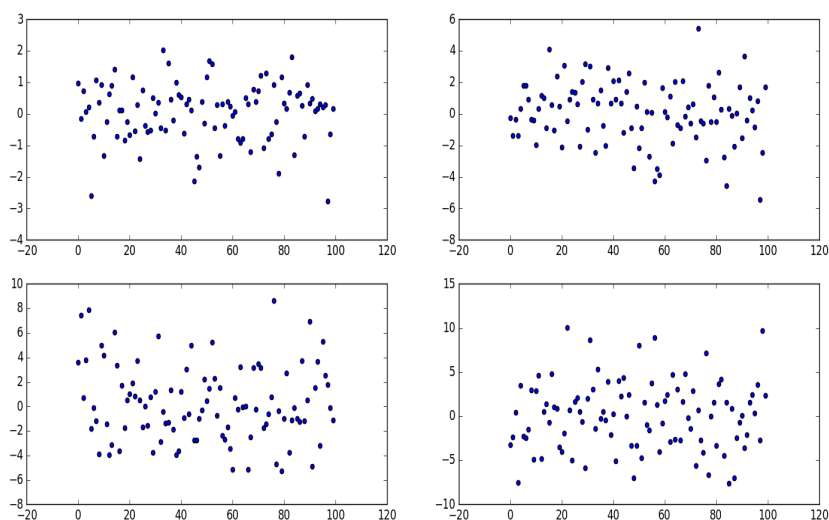


图 12.6: 题 6

```

7. In [1]: %paste
1 tan_value = np.tan(np.linspace(0,2*np.pi,10001))[np.newaxis,:]
2 tan_value=np.concatenate((np.linspace(0,2*np.pi,10001)[np.newaxis,:],
3     tan_value),0)
4 tan_value=np.where(tan_value>200,200,
5     np.where(tan_value<-200,-200,tan_value))
6 plt.plot(tan_value[0],tan_value[1],'o')
7 plt.show()
8
9
10 ## — End pasted text —

```

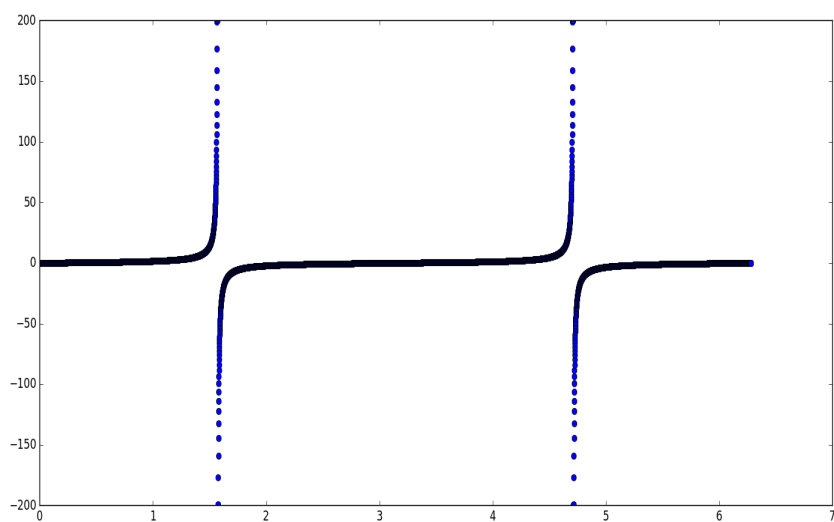


图 12.7: 题 7

第二部分

统计

第十三章 描述性统计

1. (a)

```

1 In [1]: import pandas as pd
2
3 In [2]: history = pd.read_csv('Data/Part2/001/history.csv',
4     ...:                       index_col = 'Date')
5
6 In [3]: history.index = pd.to_datetime(history.index, format='%Y-%m-%d')
7
8 In [4]: history.head()
9 Out[4]:
10
11      Convertible.Arbitrage  CTA.Global  Distressed.Securities \
12 Date
13 1997-01-31              0.0119      0.0393              0.0178
14 1997-02-28              0.0123      0.0298              0.0122
15 1997-03-31              0.0078     -0.0021             -0.0012
16 1997-04-30              0.0086     -0.0170              0.0030
17 1997-05-31              0.0156     -0.0015              0.0233
18
19      Emerging.Markets  Equity.Market.Neutral  Event.Driven \
20 Date
21 1997-01-31          0.0791              0.0189          0.0213
22 1997-02-28          0.0525              0.0101          0.0084
23 1997-03-31         -0.0120              0.0016         -0.0023
24 1997-04-30          0.0119              0.0119         -0.0005
25 1997-05-31          0.0315              0.0189          0.0346
26
27      Fixed.Income.Arbitrage  Global.Macro  Long.Short.Equity \
28 Date
29 1997-01-31          0.0191              0.0573          0.0281
30 1997-02-28          0.0122              0.0175         -0.0006
31 1997-03-31          0.0109             -0.0119         -0.0084
32 1997-04-30          0.0130              0.0172          0.0084
33 1997-05-31          0.0118              0.0108          0.0394
34
35      Merger.Arbitrage  Relative.Value  Short.Selling
36 Date

```

33	1997-01-31	0.0150	0.0180	-0.0166
34	1997-02-28	0.0034	0.0118	0.0426
35	1997-03-31	0.0060	0.0010	0.0778
36	1997-04-30	-0.0001	0.0122	-0.0129
37	1997-05-31	0.0197	0.0173	-0.0737

```
1 In [5]: history[ Emerging.Market$ ].mean()
2 Out[5]: 0.0082460526315789474
```

(b)

```
1 In [6]: history[ Emerging.Market$ ].median()
2 Out[6]: 0.013649999999999999
```

```
1 In [7]: history[ Emerging.Market$ ].mode()
2 Out[7]:
3 0      0.016
4 dtype: float64
```

(d)

```
1 In [8]: history[ Emerging.Market$ ].quantile([0.1,0.9])
2 Out[8]:
3 0.1      -0.03844
4 0.9       0.04798
5 Name: Emerging.Market$, dtype: float64
```

2.

(a)

```
1 In [9]: history[ Event.Driver$ ].max() - history[ Event.Driver$ ].min()
2 Out[9]: 0.1328
```

```
1 In [10]: history[ Event.Driver$ ].mad()
2 Out[10]: 0.012804657202216066
```

(b)

```
1 In [11]: history[ Event.Driver$ ].var()
2 Out[11]: 0.00033673989369118157
3
4 In [12]: history[ Event.Driver$ ].std()
5 Out[12]: 0.018350473936418688
```

```
1 In [13]: history[[' Relative.Value' , Fixed.Income.Arbitrage' ]].plot()
2 Out[13]: <matplotlib.axes._subplots.AxesSubplot at 0x7f3eef48a908>
```

可以看到，两者的收益率水平非常接近。为了验证这个结论，我们可以求出两者的平均数与标准差。

```
1 In [14]: history[ 'Relative.Value' ].mean()
2 Out[14]: 0.0067013157894736854
3
4 In [15]: history[ 'Fixed.Income.Arbitrage' ].mean()
5 Out[15]: 0.0042309210526315791
6
7 In [16]: history[ 'Relative.Value' ].std()
8 Out[16]: 0.0131946807807631
9
10 In [17]: history[ 'Fixed.Income.Arbitrage' ].std()
11 Out[17]: 0.014171294713188018
```

显然，两者的平均数与标准差都非常接近。因此，我们的结论是正确的。



图 13.1: 题 3

```
3.
1 In [18]: history.describe()
2 Out[18]:
3      Convertible.Arbitrage  CTA.Global  Distressed.Securities \
4 count          152.000000    152.000000          152.000000
5 mean              0.006409      0.006489              0.007953
6 std              0.020047      0.025131              0.018348
7 min             -0.123700     -0.054300             -0.083600
8 25%               0.000925     -0.011500             -0.000175
9 50%               0.009200      0.005250              0.009700
10 75%              0.014600      0.022675              0.018225
```

11	max	0.061100	0.069100	0.050400
12				
13		Emerging . Markets	Equity . Market . Neutral	Event . Driven \
14	count	152.000000	152.000000	152.000000
15	mean	0.008246	0.006003	0.007622
16	std	0.038571	0.009006	0.018350
17	min	-0.192200	-0.058700	-0.088600
18	25%	-0.011400	0.002400	-0.000025
19	50%	0.013650	0.006300	0.010150
20	75%	0.028875	0.010125	0.018650
21	max	0.123000	0.025300	0.044200
22				
23		Fixed . Income . Arbitrage	Global . Macro	Long . Short . Equity \
24	count	152.000000	152.000000	152.000000
25	mean	0.004231	0.007672	0.007760
26	std	0.014171	0.017020	0.022174
27	min	-0.086700	-0.031300	-0.067500
28	25%	0.002275	-0.003600	-0.003700
29	50%	0.006000	0.006200	0.010150
30	75%	0.009950	0.016450	0.021450
31	max	0.036500	0.073800	0.074500
32				
33		Merger . Arbitrage	Relative . Value	Short . Selling
34	count	152.000000	152.000000	152.000000
35	mean	0.006785	0.006701	0.004161
36	std	0.011168	0.013195	0.055099
37	min	-0.054400	-0.069200	-0.134000
38	25%	0.001775	0.001850	-0.025450
39	50%	0.007700	0.008450	-0.000200
40	75%	0.013725	0.014500	0.035925
41	max	0.027200	0.039200	0.246300

第十四章 随机变量

```
1.
1 In [1]: import pandas as pd
2
3 In [2]: Bwages = pd.read_csv(' Data/Part2/002/Bwages.csv ')
```

(a)

```
1 In [3]: Bwages.wage.hist(normed=True)
2 Out[3]: <matplotlib.axes._subplots.AxesSubplot at 0x7f8c7c879fd0>
```

```
1 In [4]: Bwages.wage.hist(normed=True,cumulative=True)
2 Out[4]: <matplotlib.axes._subplots.AxesSubplot at 0x7f8c7c7f75f8>
```

(b)

```
1 In [5]: from scipy import stats
2
3 In [6]: import matplotlib.pyplot as plt
4
5 In [7]: import numpy as np
6
7 In [8]: kde = stats.gaussian_kde(Bwages.wage)
8
9 In [9]: bins=np.linspace(0, 50, num=200)
10
11 In [10]: plt.plot(bins,kde(bins).cumsum())
12 Out[10]: [<matplotlib.lines.Line2D at 0x7f8c7c01fe80>]
```

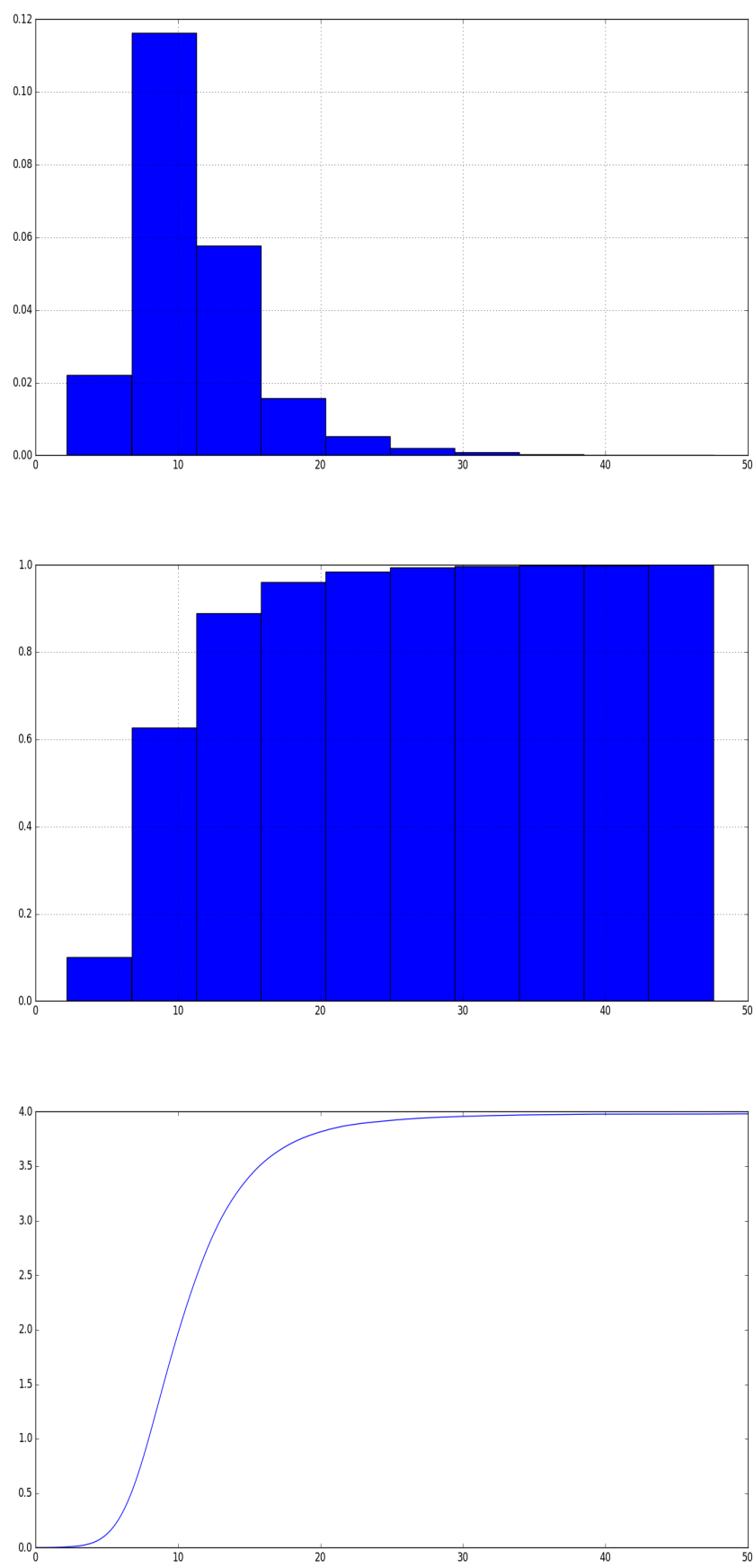


图 14.1: 题 1

```
1 In [11]: history = pd.read_csv(' Data/Part2/001/history.csv' ,index_col = ' Date' )
```

(a)

```
1 In [12]: revenue=len(history[ Emerging.Market$ ][history[ Emerging.Market$ ]>0])
```

```
1 In [13]: loss=len(history[ Emerging.Market$ ][history[ Emerging.Market$ ]<0])
```

(b)

```
1 In [14]: p=revenue/(revenue+loss)
```

```
1 In [15]: 1-stats.binom.cdf(6,12,p)
```

```
2 Out[15]: 0.86700957009568247
```

```
1 In [16]: from math import sqrt
2
3 In [17]: norm_bins=np.linspace(-5, 5, num=200)
4
5 In [18]: plt.plot(norm_bins, stats.norm.pdf(norm_bins,0,1) ,label=' N(0,1)' )
6     ...:
7     ...: plt.plot(norm_bins, stats.norm.pdf(norm_bins,0,sqrt(0.5)) ,
8     ...:           label=' N(0,0.5)' )
9     ...:
10    ...: plt.plot(norm_bins, stats.norm.pdf(norm_bins,0,sqrt(2)) ,label=' N(0,2)' )
11    ...:
12    ...: plt.plot(norm_bins, stats.norm.pdf(norm_bins,2,1) ,label=' N(2,1)' )
13    ...:
14    ...: plt.legend()
15 Out[18]: <matplotlib.legend.Legend at 0x7f8c6c47bb38>
```

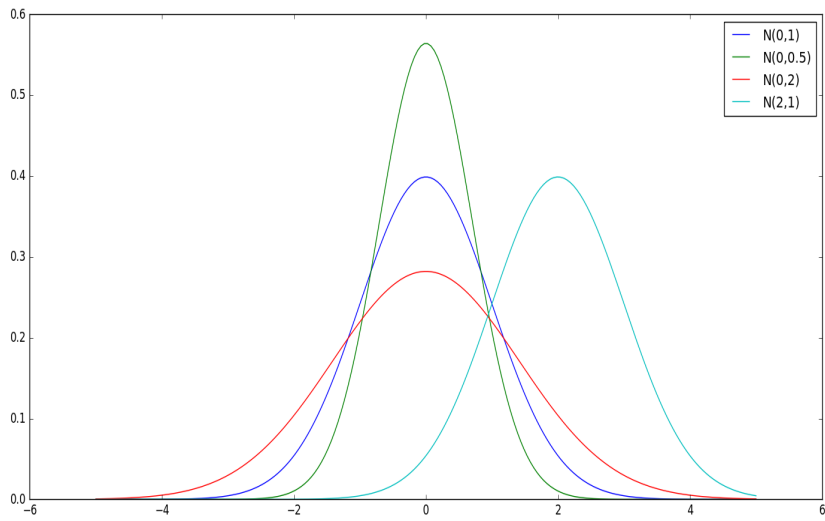


图 14.2: 题 3

(d)

(a)

$$P(x) = C_{x-1}^{r-1} p^r (1-p)^{x-r}$$

(b)

$$\begin{cases} E(x) = \frac{1}{p} \\ Var(x) = \frac{1-p}{p^2} \end{cases}$$

5.

(a)

$$F(x) = -exp\left(-\frac{x}{\lambda}\right)$$

(b)

$$\begin{cases} E(x) = \lambda \\ Var(x) = \lambda^2 \end{cases}$$

(c)

```
1 In [19]: import numpy as np
2
3 In [20]: sample = np.random.exponential(2,10000)
4
5 In [21]: sample.mean()
6 Out[21]: 2.0077693716605811
7
8 In [22]: sample.var()
```

```
9 Out[22]: 3.9950035039304801
```

6.

(a)

$$\begin{cases} E(x) = \exp\left(\frac{\mu + \sigma^2}{2}\right) \\ \text{Var}(x) = (\exp(\sigma^2) - 1) \exp(2\mu + \sigma^2) \end{cases}$$

(b) 假设均值和标准差均为 1

```
1 In [23]: import matplotlib.pyplot as plt
2
3 In [24]: log_bins=np.linspace(0, 10, num=200)
4
5 In [25]: plt.plot(log_bins, stats.lognorm.pdf(log_bins,1,0,1))
6 Out[25]: [<matplotlib.lines.Line2D at 0x7f8c6c403c50>]
```

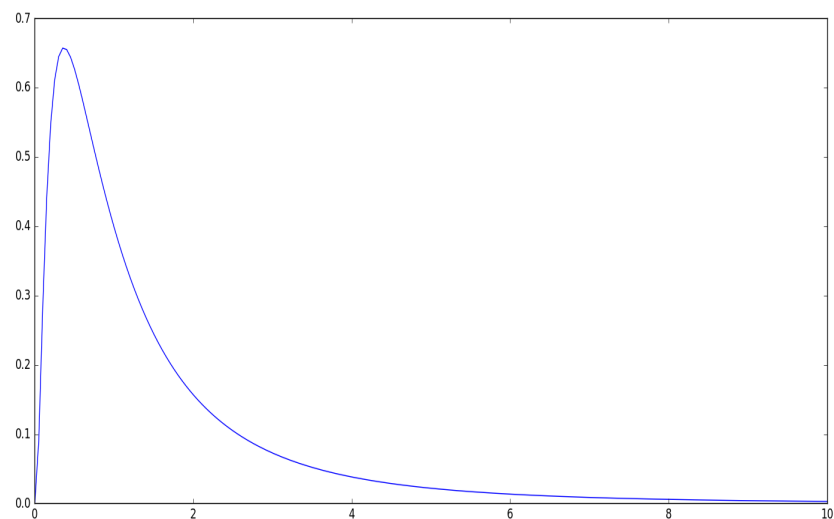


图 14.3: 题 6

第十五章 推断统计

1.

(a)

```
1 In [1]: import numpy as np
2
3 In [2]: import math
4
5 In [3]: from scipy import stats
6
7 In [4]: mu=np.mean((6.683,6.678,6.767,6.692,6.672,6.678))
8
9 In [5]: std=np.std((6.683,6.678,6.767,6.692,6.672,6.678))
10
11 In [6]: low = mu - stats.t.ppf(0.95,5) * std / math.sqrt(6)
12
13 In [7]: high = mu + stats.t.ppf(0.95,5) * std / math.sqrt(6)
14
15 In [8]: low
16 Out[8]: 6.6680404901748709
17
18 In [9]: high
19 Out[9]: 6.7219595098251279
```

$\therefore \mu$ 的置信水平为 0.9 的置信区间为 (6.668, 6.722)

(b)

```
1 In [9]: mu=np.mean((6.661,6.664,6.668,6.666,6.665))
2
3 In [10]: std=np.std((6.661,6.664,6.668,6.666,6.665))
4
5 In [11]: low = mu - stats.t.ppf(0.95,4) * std / math.sqrt(5)
6
7 In [12]: high = mu + stats.t.ppf(0.95,4) * std / math.sqrt(5)
8
9 In [13]: low
10 Out[13]: 6.6625927405704291
```

```

11
12 In [14]: high
13 Out[14]: 6.6670072594295702

```

$\therefore \mu$ 的置信水平为 0.9 的置信区间为 (6.663, 6.667)

```

1 In [15]: mu=np.mean((5.9,7.3,6.6,5.8,5.7,5.3,5.9,7,6.5))
2
3 In [16]: std=np.std((5.9,7.3,6.6,5.8,5.7,5.3,5.9,7,6.5))
4
5 In [17]: low = mu - stats.t.ppf(0.975,8) * std / 3
6
7 In [18]: high = mu + stats.t.ppf(0.975,8) * std / 3
8
9 In [19]: low
10 Out[19]: 5.7431779225460415
11
12 In [20]: high
13 Out[20]: 6.7012665218984031

```

$\therefore \mu$ 的置信水平为 0.95 的置信区间为 (5.743, 6.701)。

2. $\because \bar{x}$ 服从 $N\left(\mu, \frac{\sigma^2}{n}\right)$

$$\therefore \begin{cases} E(\bar{x}) = \mu \\ Var(\bar{x}) = \frac{\sigma^2}{n} \end{cases}$$

$$\therefore \begin{cases} E\left(\frac{\bar{x}-\mu}{\sigma/\sqrt{n}}\right) = \frac{E(\bar{x})-\mu}{\sigma/\sqrt{n}} = 0 \\ Var\left(\frac{\bar{x}-\mu}{\sigma/\sqrt{n}}\right) = \frac{Var(\bar{x}-\mu)}{\sigma^2/n} = \frac{\sigma^2/n}{\sigma^2/n} = 1 \end{cases}$$

又因为服从正态分布的随机变量和一个常数进行加减乘除运算后仍服从正态分布，所以 $\frac{\bar{x}-\mu}{\sigma/\sqrt{n}} \sim N(0, 1)$

4. 一阶矩估计量为样本均值: $\bar{x} = \frac{x_1+x_2+\dots+x_n}{n}$;

二阶矩估计量为样本方差: $s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$

```

5.
1 In [1]: import pandas as pd
2
3 In [2]: import matplotlib.pyplot as plt
4
5 In [3]: Bwages = pd.read_csv(' Bwages.csv ')
6
7 In [4]: mu = Bwages.wage.mean()
8
9 In [5]: std=Bwages.wage.std()

```

```
10
11 In [6]: low = mu - stats.t.ppf(0.975, len(Bwages)-1) * std / math.sqrt(len(Bwages)
12         )
13 In [7]: high = mu + stats.t.ppf(0.975, len(Bwages)-1) * std / math.sqrt(len(Bwages)
14         ))
15 In [8]: low
16 Out[8]: 10.823074026756808
17
18 In [9]: high
19 Out[9]: 11.278158191999699
```

∴wage 的置信水平为 0.95 的置信区间为 (10.82307, 11.27816)。

```
6.
1 In [10]: Bwages.wage.hist(normed=True)
2 Out[10]: <matplotlib.axes._subplots.AxesSubplot at 0x7fa538a0a940>
3
4 In [11]: bins=np.linspace(0,50,200)
5
6 In [12]: plt.plot(bins, stats.norm.pdf(bins, mu, std))
7 Out[12]: [<matplotlib.lines.Line2D at 0x7fa53808cb38>]
```

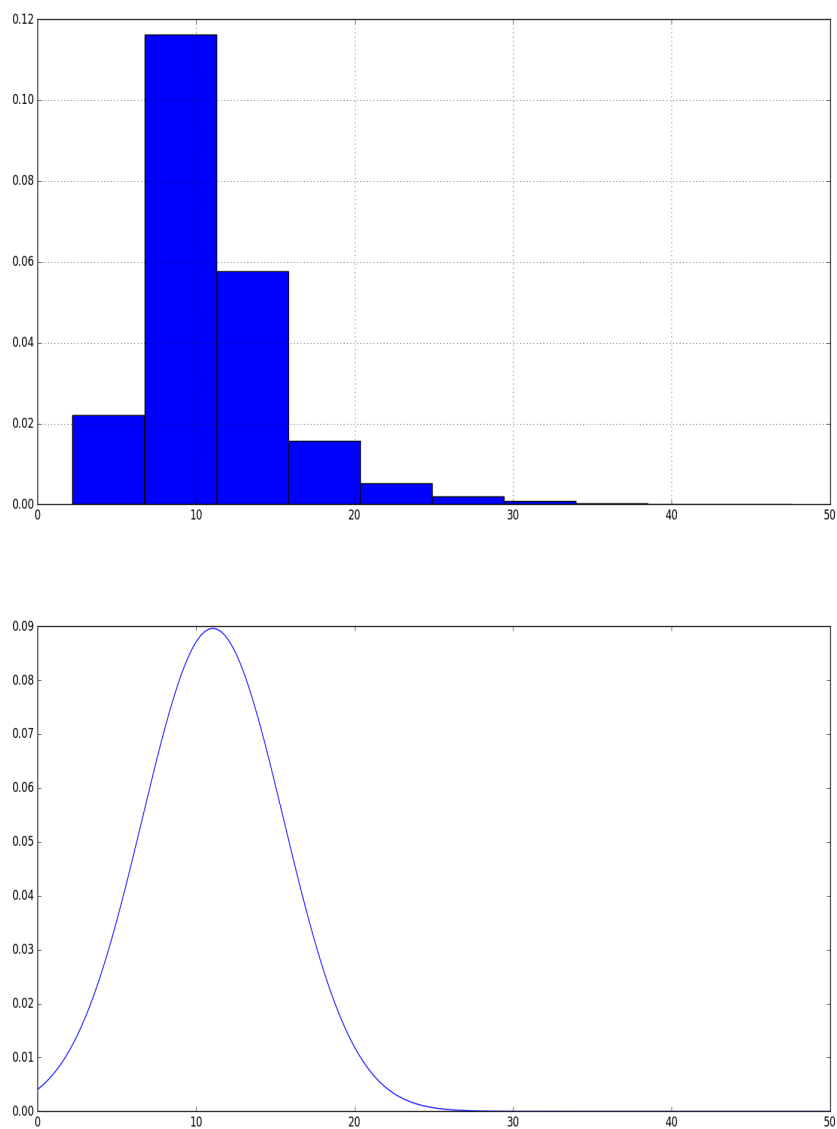


图 15.1: 题 6

可以看到, `wage` 的频率直方图与正态分布的曲线是较为接近的, 所以我们的假设是合理的。

```
7. In [1]: from scipy import stats
8. In [2]: MT=(0.225,0.262,0.217,0.24,0.23,0.229,0.235,0.217)
9. In [3]: Sn=(0.209,0.205,0.196,0.21,0.202,0.207,0.224,0.223)
10. In [4]: stats.ttest_ind(MT,Sn)
11. Out[4]: Ttest_indResult(statistic=3.6170843653195446,
12.         pvalue=0.0028018899619853565)
```

因为 p 值为 0.002802, 小于 0.05, 所以我们应当拒绝原假设。即, 两位作家所写的小品文中包含由 3 个字母组成的单字的比例具有显著性的差异。

```

8.
1 In [1]: from scipy import stats
2
3 In [2]: region1=(168,180,181,172,165,160,166,165,177,174)
4
5 In [3]: region2=(169,170,176,173,166,167,166,173,171,170)
6
7 In [4]: stats.ttest_ind(region1,region2)
8 Out[4]: Ttest_indResult(statistic=0.28303531904272478,
9         pvalue=0.78037906653623068)

```

因为 p 值为 0.7804，大于 0.05，所以我们不能拒绝原假设。因此，这两个地区的人的身高不具有显著性差异。

```

9.
1 In [1]: import pandas as pd
2
3 In [2]: from scipy import stats
4
5 In [3]: Bwages = pd.read_csv(' Data/Part2/002/Bwages.csv ')
6
7 In [4]: stats.ttest_1samp(Bwages.wage,11)
8 Out[4]: Ttest_1sampResult(statistic=0.4363476178926709, pvalue
          =0.66264857119758025)

```

因为 p 值为 0.6626，大于 0.05，所以我们不能拒绝原假设。

```

10.
1 In [1]: import pandas as pd
2
3 In [2]: from scipy import stats
4
5 In [3]: history = pd.read_csv(' Data/Part2/001/history.csv ',
6         ....:                  index_col = ' Date ')
7
8 In [4]: stats.ttest_ind(history[' Emerging.Markets' ],
9         ....:            history[' Global.Macro '])
10 Out[4]: Ttest_indResult(statistic=0.1677641921908746,
11         pvalue=0.86688108106493611)

```

因为 p 值为 0.8669，大于 0.05，所以我们不能拒绝原假设。检验新兴市场风格对冲基金和全球宏观风格对冲基金收益率不具有显著性差异。

```

11.
1 In [5]: stats.ttest_rel(history[' Emerging.Markets' ],
2         ....:            history[' Global.Macro '])
3 Out[5]: Ttest_relResult(statistic=0.23846933281992622,

```

```
4 pvalue=0.81184041604248691)
```

因为 p 值为 0.8118，大于 0.05，所以我们不能拒绝原假设。

第十六章 方差分析

1.

$$\begin{aligned} FSS &= \sum_{i=1}^5 10 \times (\bar{x}_i - \bar{x})^2 \\ &= 10 \times (11.4 - 10.7)^2 + 10 \times (11.7 - 10.7)^2 \\ &+ 10 \times (9.9 - 10.7)^2 + 10 \times (9.8 - 10.7)^2 + 10 \times (7.9 - 10.7)^2 \\ &= 107.8 \end{aligned}$$

$$ESS = \sum_{j=1}^5 \sum_{i=1}^{10} (x_{ij} - \bar{x}_j)^2 = 445.9$$

$$F = \frac{FSS/4}{ESS/45} = 2.72$$

我们可以查到 $F_{4,45,0.95}$ 的值为 2.56。因为 F 统计量的值大于 2.56，所以我们应当拒绝原假设。即，5 个地区的单日开户数量具有显著性差异。

2.

$$\begin{aligned} FSS &= \sum_{i=1}^4 5 \times (\bar{x}_i - \bar{x})^2 \\ &= 5 \times (1.138\% - 0.5785\%)^2 + 5 \times (0.002\% - 0.5785\%)^2 \\ &+ 5 \times (0.22\% - 0.5785\%)^2 + 5 \times (0.954\% - 0.5785\%)^2 \\ &= 0.046\% \end{aligned}$$

$$ESS = \sum_{j=1}^4 \sum_{i=1}^5 (x_{ij} - \bar{x}_j)^2 = 0.42\%$$

$$F = \frac{FSS/3}{ESS/16} = 0.58$$

我们可以查到 $F_{3,16,0.95}$ 的值为 3.24。因为 F 统计量的值小于 3.24，所以我们不能拒绝原假设。即，调息方案对上证综指不具有显著性影响。

3.

$$\begin{aligned}
 FSS &= \sum_{i=1}^3 5 \times (\bar{x}_i - \bar{x})^2 \\
 &= 5 \times (43.6 - 39.87)^2 + 5 \times (30.2 - 39.87)^2 + 5 \times (45.8 - 39.87)^2 \\
 &= 712.93
 \end{aligned}$$

$$ESS = \sum_{j=1}^3 \sum_{i=1}^5 (x_{ij} - \bar{x}_j)^2 = 278.8$$

$$F = \frac{FSS/2}{ESS/12} = 15.34$$

我们可以查到 $F_{2,12,0.95}$ 的值为 3.89。因为 F 统计量的值大于 3.89，所以我们应当拒绝原假设。即，电池的平均寿命具有显著性的差异。

4.

```

1 In [1]: import pandas as pd
2
3 In [2]: managers = pd.read_csv('Data/Part2/004/managers.csv', index_col=0)
4
5 In [3]: MANA = managers.iloc[:, (0, 2, 3)]

```

(a)

```

1 In [4]: ((MANA.HAM1-MANA.HAM1.mean())**2).sum()
2 Out[4]: 0.08604549181818183
3
4 In [5]: ((MANA.HAM3-MANA.HAM3.mean())**2).sum()
5 Out[5]: 0.1746451887878788
6
7 In [6]: ((MANA.HAM4-MANA.HAM4.mean())**2).sum()
8 Out[6]: 0.37073304333333335

```

```

1 In [7]: mu = MANA.mean().mean()
2
3 In [8]: ((MANA.HAM1.mean()-mu)**2 + (MANA.HAM3.mean()-mu)**2 + (MANA.HAM4.mean()
4 Out[8]: 0.00016766787878787911

```

(b)

```

1 In [9]: 0.0001676679+0.08604549+0.1746452+0.370733
2 Out[9]: 0.6315913579

```

(d) ESS、FSS、TSS 的自由度分别为 393、2、395。三者的关系为： $df_{ESS} + df_{FSS} = df_{TSS}$ 。

(e) $F = \frac{0.0002/2}{(0.086+0.175+0.371)/393} = 0.062$

我们可以查到 $F_{2,12,0.95}$ 的值为 3.07。因为 F 统计量的值小于 3.07，所以我们不能拒绝原假设。即，这三者的平均收益率不具有显著性差异。

```

1 In [10]: from statsmodels.formula.api import ols
2
3 In [11]: import statsmodels.stats.anova as anova
4
5 In [12]: returns = pd.DataFrame(pd.concat([MANA.HAM1,MANA.HAM3,MANA.HAM4]))
6
7 In [13]: returns['Class'] = ['HAM1' for i in range(132)]+['HAM3' for i in range
8                               (132)]+['HAM4' for i in range(132)]
9
10 In [14]: returns.columns = ['Return' , 'Class']
11
12 In [15]: model = ols('Return~C(Class)',data=returns).fit()
13
14 In [16]: print(anova.anova_lm(model))
15
16      df    sum_sq  mean_sq      F    PR(>F)
17 C(Class)    2.0  0.000168  0.000084  0.052178  0.949166
18 Residual  393.0  0.631424  0.001607      NaN      NaN

```

可以看到，`anova.anova_lm()` 得到的 F 取值与我们自己计算得到的结果很相近。

第十七章 回归

5.

(a)

```
1 In [1]: import matplotlib.pyplot as plt
2
3 In [2]: x = list(range(1952,2016,4))
4
5 In [3]: y = (29.3,28.8,28.5,28.4,29.4,27.6,27.7,27.7,
6           ...:      27.8,27.4,27.8,27.1,27.3,27.1,27.0,27.5)
7
8 In [4]: plt.plot(x,y)
9 Out[4]: [<matplotlib.lines.Line2D at 0x7f4eec296518>]
```

$$(b) \bar{x} = \frac{\sum_{i=1}^{16} x_i}{16} = 1982$$

$$\bar{y} = \frac{\sum_{i=1}^{16} y_i}{16} = 27.9$$

$$\hat{b} = \frac{\sum_{i=1}^{16} (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^{16} (x_i - \bar{x})^2} = -0.034$$

$$\hat{a} = \bar{y} - \hat{b}\bar{x} = 94.647$$

y 关于 x 的线性回归方程为 $y = 94.647 - 0.034x$

```
1 In [5]: import statsmodels.api as sm
2
3 In [6]: model=sm.OLS(y ,sm.add_constant(x)).fit()
4
5 In [7]: print(model.summary())
```

OLS Regression Results						
=====						
Dep. Variable:	y	R-squared:	0.711			
Model:	OLS	Adj. R-squared:	0.690			
Method:	Least Squares	F-statistic:	34.41			
Date:	Fri, 20 Jan 2017	Prob (F-statistic):	4.11e-05			
Time:	18:20:40	Log-Likelihood:	-7.8861			
No. Observations:	16	AIC:	19.77			
Df Residuals:	14	BIC:	21.32			
Df Model:	1					
Covariance Type:	nonrobust					
=====						
	coef	std err	t	P> t	[95.0% Conf. Int.]	

const	94.6468	11.380	8.317	0.000	70.239	119.054
x1	-0.0337	0.006	-5.866	0.000	-0.046	-0.021
=====						
Omnibus:	5.698	Durbin-Watson:	2.210			
Prob(Omnibus):	0.058	Jarque-Bera (JB):	3.021			
Skew:	0.985	Prob(JB):	0.221			
Kurtosis:	3.807	Cond. No.	2.13e+05			
=====						

(c) $t = \frac{\hat{b}\sqrt{\sum_{i=1}^{16}(x_i-\bar{x})^2}}{\hat{\sigma}} = -5.866$

因为 $t_{14,0.975} = 2.145 < |t|$ ，所以我们可以拒绝原假设。

(d) 将 $x = 2016$ 代入可得： $y = 26.103$ 。所以，2016 年里约奥运会男子 10000 米冠军的成绩的预测值为 26.103。

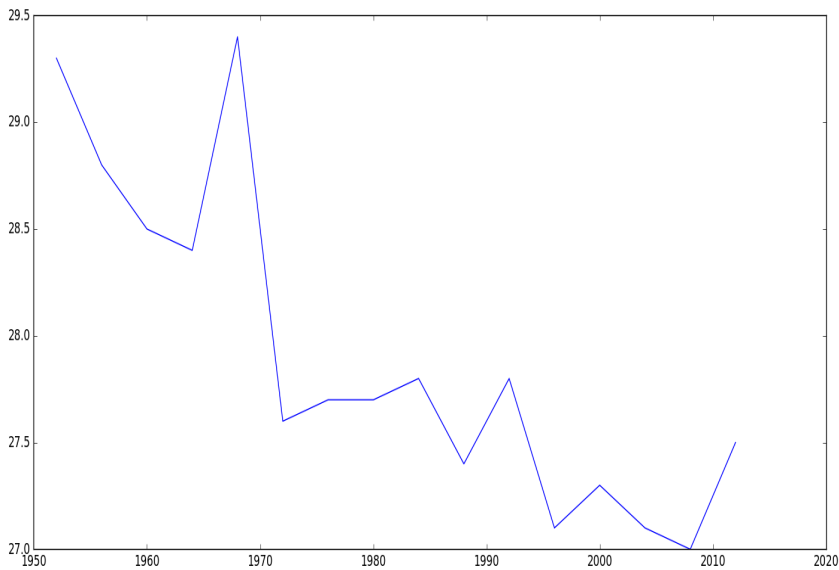


图 17.1: 题 1

```
2. In [1]: import pandas as pd
3. In [2]: EU = pd.read_csv(' Data/Part2/005/EuStockMarkets.csv ')
4.
5. In [3]: plt.plot(EU.DAX,EU.FTSE,' . ' )
6. ...:
7. ...: plt.xlabel(' DAX ' )
```



```
(b)
1 In [4]: plt.plot(EU.FTSE,EU.DAX,'.',EU.FTSE,model.fittedvalues,'-')
2     ...:
3     ...: plt.xlabel(' FTSE ')
4     ...:
5     ...: plt.ylabel(' DAX ')
6 Out[4]: <matplotlib.text.Text at 0x7f4ed2e6e4e0>
```

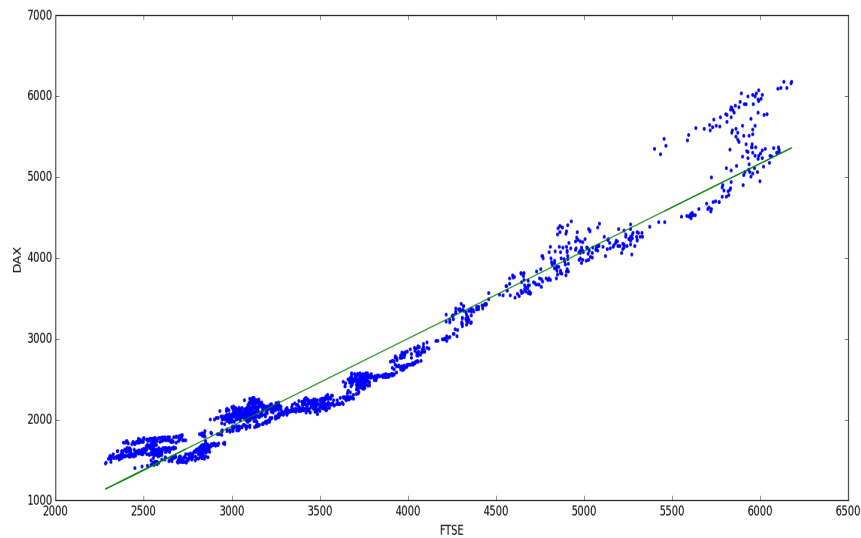


图 17.3: 题 3

```
4.
1 In [1]: plt.plot(model.fittedvalues,model.resid,'.')
2     ...:
3     ...: plt.xlabel(' Fitted ')
4     ...:
5     ...: plt.ylabel(' Residual ')
6 Out[1]: <matplotlib.text.Text at 0x7f4ed2e58748>
7
8 In [2]: import scipy.stats as stats
9
10 In [3]: sm.qqplot(model.resid_pearson,stats.norm,line=45)
11
12 In [4]: plt.plot(model.fittedvalues,model.resid_pearson**0.5,'.')
13     ...:
14     ...: plt.xlabel(' Fitted ')
15     ...:
16     ...: plt.ylabel(' Square Root of Standardized Residual ')

```

在第一张图和第三张图中，我们都可以看到很明显的趋势。这说明线性及同方差的条件并没有被满足。而第二张图也存在严重的偏离，表明因变量并不服从正态分布。因此，我们的模型并不能很好地模拟两者之间的关系。实际上，DAX 指数与 FTSE 指数属于时间序列，我们应该使用时间序列模型而不是线性回归模型。

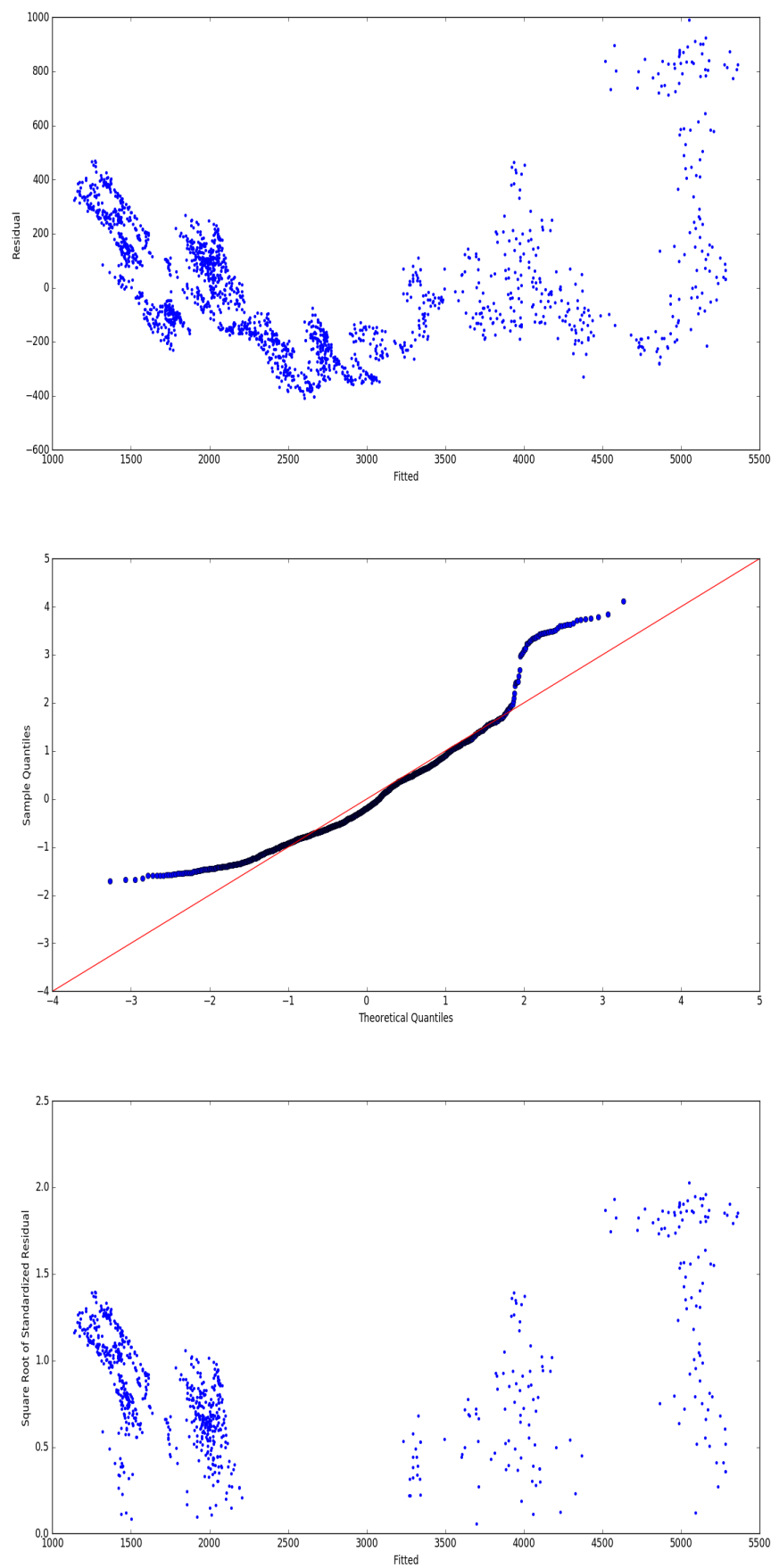


图 17.4: 题 4

5.

(a)

```

1 In [1]: import matplotlib.pyplot as plt
2         ...: import statsmodels.api as sm
3         ...: import numpy as np
4         ...:
5
6 In [2]: x = [20,25,30,35,40,50,60,65,70,75,80,90]
7
8 In [3]: y = [1.81,1.7,1.65,1.55,1.48,1.4,1.3,1.26,1.24,1.21,1.2,1.18]
9
10 In [4]: plt.plot(x,y,' ')
11 Out[4]: [<matplotlib.lines.Line2D at 0x7f4ed2cddb70>]

```

```

1 In [5]: independent = np.array([x,[i**2 for i in x]]).T
2
3 In [6]: model = sm.OLS(y,sm.add_constant(independent)).fit()
4
5 In [7]: print(model.summary())

```

```

=====
                        OLS Regression Results
=====
Dep. Variable:          y      R-squared:                0.997
Model:                  OLS    Adj. R-squared:            0.997
Method:                 Least Squares    F-statistic:        1767.
Date:                   Tue, 17 Jan 2017    Prob (F-statistic):    2.10e-12
Time:                   20:01:19    Log-Likelihood:        37.586
No. Observations:       12    AIC:                   -69.17
Df Residuals:           9    BIC:                   -67.72
Df Model:                2
Covariance Type:        nonrobust
=====

```

	coef	std err	t	P> t	[95.0% Conf. Int.]	
const	2.1983	0.023	97.485	0.000	2.147	2.249
x1	-0.0225	0.001	-23.898	0.000	-0.025	-0.020
x2	0.0001	8.66e-06	14.446	0.000	0.000	0.000

```

=====
Omnibus:                 0.874    Durbin-Watson:           2.783
Prob(Omnibus):            0.646    Jarque-Bera (JB):         0.642
Skew:                    -0.096    Prob(JB):                 0.726
Kurtosis:                 1.884    Cond. No.:                2.65e+04
=====

```

(b) 根据回归结果中的 F Statistic，我们可以拒绝原假设。

(d)

```

1 In [8]: model.predict(np.array([1,95,95**2]).T)
2 Out[8]: array([ 1.18735552])

```

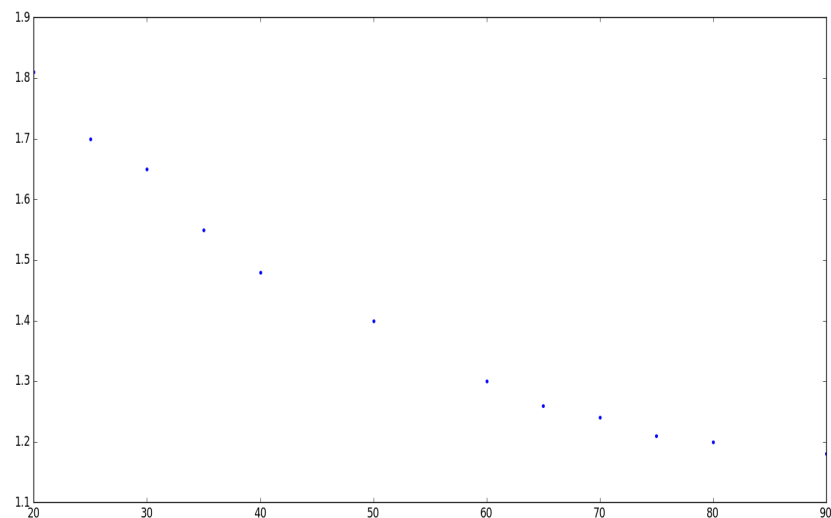


图 17.5: 题 4

```

6. In [1]: import pandas as pd
    In [2]: import statsmodels.formula.api as smf
    In [3]: import numpy as np
    In [4]: cps = pd.read_csv('Data/Part2/005/CPS1988.csv')

```

```

(a) In [5]: cps.head()
Out[5]:
   wage  education  experience  ethnicity  smsa  region  parttime
0  354.94         7         45      cauc   yes  northeast      no
1  123.46        12          1      cauc   yes  northeast     yes
2  370.37         9          9      cauc   yes  northeast      no
3  754.94        11         46      cauc   yes  northeast      no
4  593.54        12         36      cauc   yes  northeast      no

```

```

In [6]: model = smf.ols(' np.log(wage)~experience+education+ethnicity ',
...:                    data=cps).fit()

```

```

(b) In [7]: print(model.summary())

```

```

OLS Regression Results

=====
Dep. Variable:      np.log(wage)    R-squared:          0.221
Model:              OLS             Adj. R-squared:     0.221
Method:             Least Squares   F-statistic:        2665.
Date:               Tue, 17 Jan 2017 Prob (F-statistic):  0.00
Time:               20:04:22         Log-Likelihood:      -27020.
No. Observations:   28155           AIC:                5.405e+04
Df Residuals:       28151           BIC:                5.408e+04
Df Model:           3
Covariance Type:    nonrobust

=====
                    coef    std err          t      P>|t|      [95.0% Conf. Int.]
-----
Intercept          4.2889      0.023    183.656    0.000      4.243      4.335
ethnicity[T.cauc]  0.2431      0.014    17.395    0.000      0.216      0.271
experience          0.0196      0.000    65.198    0.000      0.019      0.020
education          0.0996      0.001    73.258    0.000      0.097      0.102

=====
Omnibus:            2846.769    Durbin-Watson:      1.814
Prob(Omnibus):      0.000    Jarque-Bera (JB):   4634.595
Skew:               -0.733    Prob(JB):           0.00
Kurtosis:           4.343    Cond. No.           163.

=====

```

(d)

```

1 In [8]: model.fittedvalues.head()
2 Out[8]:
3 0    6.110772
4 1    5.746557
5 2    5.604576
6 3    6.528665
7 4    6.432310
8 dtype: float64

```

```

1 In [9]: model2 = smf.ols(' np.log(wage)~experience+np.power(experience,2)+
2          education+ethnicity ,
          data=cps).fit()

```

7. (a)

```

1 In [10]: print(model2.summary())

```

```

OLS Regression Results

=====
Dep. Variable:      np.log(wage)    R-squared:          0.335
Model:              OLS             Adj. R-squared:     0.335
Method:             Least Squares   F-statistic:        3541.
Date:               Tue, 17 Jan 2017 Prob (F-statistic):  0.00
Time:               20:06:23         Log-Likelihood:      -24801.
No. Observations:   28155           AIC:                4.961e+04
Df Residuals:       28150           BIC:                4.965e+04
Df Model:           4
Covariance Type:    nonrobust

=====
                    coef    std err          t      P>|t|      [95.0% Conf. Int.]
-----
Intercept          4.0780      0.022    187.086    0.000      4.035      4.121
ethnicity[T.cauc]  0.2434      0.013    18.839    0.000      0.218      0.269
experience          0.0775      0.001    88.033    0.000      0.076      0.079
np.power(experience, 2) -0.0013    1.9e-05   -69.312    0.000     -0.001     -0.001
education          0.0857      0.001    67.343    0.000      0.083      0.088

=====
Omnibus:            2034.828    Durbin-Watson:      1.795
Prob(Omnibus):      0.000    Jarque-Bera (JB):   3963.309
Skew:               -0.508    Prob(JB):           0.00
Kurtosis:           4.532    Cond. No.           5.23e+03

=====

```

(b) 首先,从单个变量上看, $experience^2$ 的系数是显著的。这说明我们应该将其加入模型。而且,加入 $experience^2$ 后,模型的 R^2 得到了提升。因此,相较于第一个模型,第二个模型要更好。

8. 从结果来看, F 统计量为 3541, p 值远小于 0.05。因此,我们应该拒绝原假设。

9. 从结果来看, β_3 的 p 值远小于 0.05。因此,我们应当拒绝原假设。也就是说, β_3 显著不为 0。

第三部分

金融

第十八章 资产收益率和风险

1.

$$(a) R_4 = \frac{P_4 - P_3}{P_3} = -4.4\%$$

$$(b) R_5(2) = \frac{P_5 - P_3}{P_3} = -7.2\%$$

$$(c) R_{12}(8) = \frac{P_{12} - P_4}{P_4} = -1.2\%$$

$$2. R_{12}(8) = \prod_{i=4}^{12} (1 + R_i) - 1$$

$$3. R_t = \frac{P_t - D - P_{t-1}}{P_{t-1}} = -8.3\%$$

$$4. (1 + R)^{365} = 1 + 0.03416$$

解得: $R = 0.0092\%$

$$5. R = 1.5\% \times \frac{12}{3} = 6\%$$

$$6. R_t = \lim_{n \rightarrow \infty} \left(1 + \frac{r_t}{n}\right)^n - 1 = \lim_{n \rightarrow \infty} \left(1 + \frac{r_t}{n}\right)^{\frac{n}{r_t} r_t} - 1 = e^{r_t} - 1$$

$$7. \ln(1 + R) = 0.06$$

$$R = 6.18\%$$

$$8. R = 0.5\% \times 12 = 6\%$$

9.

```
1 In [1]: import pandas_datareader.data as web
```

```
2
```

```
3 In [2]: import datetime as dt
```

```
4
```

```
5 In [3]: import numpy as np
```

(a)

```
1 In [4]: start = dt.datetime(2014,1,1)
```

```
2
```

```
3 In [5]: end = dt.datetime(2014,12,31)
```

```
4
```

```
5 In [6]: baidu = web.DataReader(' BIDU ', yahoo , start , end)
```

```
1 In [7]: returns = (baidu.Close - baidu.Close.shift(1))/baidu.Close.shift(1)
```

(b)

```
1 In [8]: comp_returns = np.log(baidu.Close/baidu.Close.shift(1))
```

```
1 In [9]: comp_returns.sum()
```

```
2 Out[9]: 0.23631272416563825
```

(d)

```
1 In [10]: returns.plot()
```

```
1 In [11]: comp_returns.plot()
```

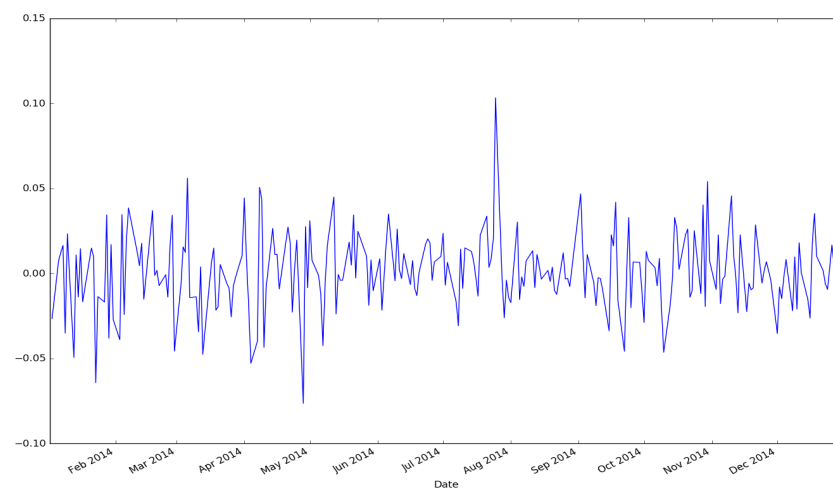
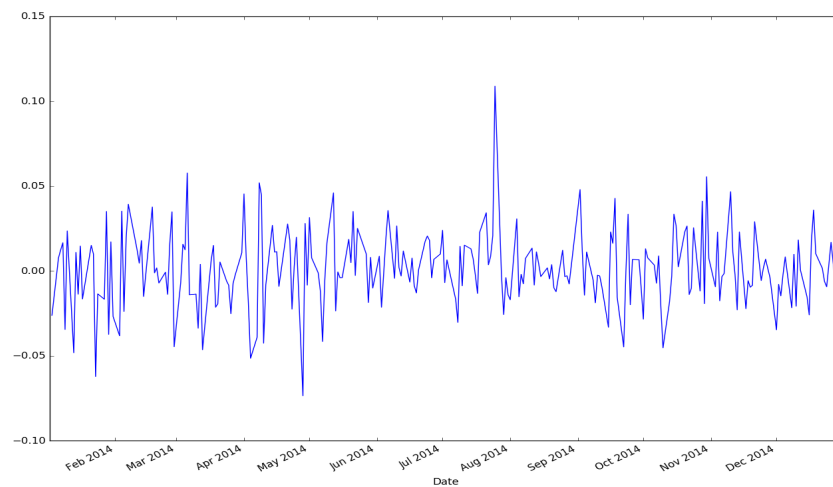


图 18.1: 题 1

(f)

10. 略

第十九章 投资组合理论及其拓展

1. 组合 A 和组合 E 的收益率都是 0.05，但是组合 E 的风险更高，因此不是有效的。同理，组合 B、C 也是无效的。组合 A 的收益和风险皆比组合 D 低，因此两者无法比较，都有可能是有效的。
2. 股票 A 和股票 D。因为两者的相关系数最小。
- 3.

$$(a) E(R) = 0.3 \times 8\% + 0.7 \times 6\% = 6.6\%$$

$$\sigma = \sqrt{(0.3 \times 0.02)^2 + (0.7 \times 0.03)^2 + 2 \times 0.3 \times 0.7 \times 0.5 \times 0.02 \times 0.03} = 0.025$$

$$(b) E(R) = 0.6 \times 8\% + 0.4 \times 6\% = 7.2\%$$

$$\sigma = \sqrt{(0.6 \times 0.02)^2 + (0.4 \times 0.03)^2 + 2 \times 0.6 \times 0.4 \times 0.5 \times 0.02 \times 0.03} = 0.021$$

4. 因为收益率为 5%，所以我们只有两种可能的组合，即 A 和 B、A 和 C。当组合为 A 和 B 时，A 和 B 的比重分别为 0.6 和 0.4，组合的标准差为 0.016。当组合为 A 和 C 时，两者的比重分别为 0.2 和 0.8，组合的标准差为 0.023。所以。我们应该选择组合 A 和 B。
5. 直接利用文中的函数，我们可以绘得最小方差前缘曲线。

```
1 In [1]: %paste
2 import numpy as np
3 import math
4 import matplotlib.pyplot as plt
5 import ffn
6 import pandas as pd
7 import pandas_datareader.data as web
8 import datetime as dt
9 from scipy import linalg
10
11 class MeanVariance:
12     def __init__(self, returns):
13         self.returns=returns
14     def minVar(self, goalRet):
15         covs=np.array(self.returns.cov())
16         means=np.array(self.returns.mean())
17         L1=np.append(np.append(covs.swapaxes(0,1),[means],0),
18                     [np.ones(len(means)),0].swapaxes(0,1)
```

```

19     L2=list(np.ones(len(means)))
20     L2.extend([0,0])
21     L3=list(means)
22     L3.extend([0,0])
23     L4=np.array([L2,L3])
24     L=np.append(L1,L4,0)
25     results=linalg.solve(L,np.append(np.zeros(len(means)),[1,goalRet],0))
26     return(np.array([list(self.returns.columns),results[: -2])))
27 def frontierCurve(self):
28     goals=[x/500000 for x in range(-100,4000)]
29     variances=list(map(lambda x: self.calVar(self.minVar(x)[1,:].astype(np.
30         float)),goals))
31     plt.plot(variances,goals)
32 def meanRet(self,fracs):
33     meanRisky=ffn.to_returns(self.returns).mean()
34     assert len(meanRisky)==len(fracs), 'Length of fractions must be equal to
35         number of assets'
36     return(np.sum(np.multiply(meanRisky,np.array(fracs))))
37 def calVar(self,fracs):
38     return(np.dot(np.dot(fracs,self.returns.cov()),fracs))

```

获取数据，绘制最小方差前缘曲线。

```

(a)
1 In [2]: %paste
2 sclq = web.DataReader(' 600039.SS ' , yahoo ,
3     dt.datetime(2009,1,2),dt.datetime(2012,12,31)).Close
4 sclqRet=((sclq-sclq.shift(1))/sclq.shift(1))
5 sclqRet.name=' sclq'
6
7 ## — End pasted text —
8
9 In [3]: %paste
10 hsly = web.DataReader(' 600054.SS ' , yahoo ,
11     dt.datetime(2009,1,2),dt.datetime(2012,12,31)).Close
12 hslyRet=((hsly-hsly.shift(1))/hsly.shift(1))
13 hslyRet.name=' hsly'
14
15 ## — End pasted text —
16
17 In [4]: %paste
18 wldc = web.DataReader(' 600173.SS ' , yahoo ,

```

```
19 dt.datetime(2009,1,2),dt.datetime(2012,12,31)).Close
20 wldcRet=((wldc-wldc.shift(1))/wldc.shift(1))
21 wldcRet.name='wldc'
22
23 ## — End pasted text —
24
25 In [5]: %paste
26 ghny = web.DataReader('600256.SS',yahoo,
27 dt.datetime(2009,1,2),dt.datetime(2012,12,31)).Close
28 ghnyRet=((ghny-ghny.shift(1))/ghny.shift(1))
29 ghnyRet.name='ghny'
30
31 ## — End pasted text —
32
33 In [6]: %paste
34 zhjt = web.DataReader('600252.SS',yahoo,
35 dt.datetime(2009,1,2),dt.datetime(2012,12,31)).Close
36 zhjtRet=((zhjt-zhjt.shift(1))/zhjt.shift(1))
37 zhjtRet.name='zhjt'
38
39 ## — End pasted text —
40
41 In [7]: %paste
42 fiveStocks=pd.concat([sclqRet,hslyRet,wldcRet,ghnyRet,zhjtRet],1)
43 fiveStocks=fiveStocks.dropna()
44
45 ## — End pasted text —
46
47 In [8]: minVar=MeanVariance(fiveStocks)
48
49 In [9]: minVar.frontierCurve()
```

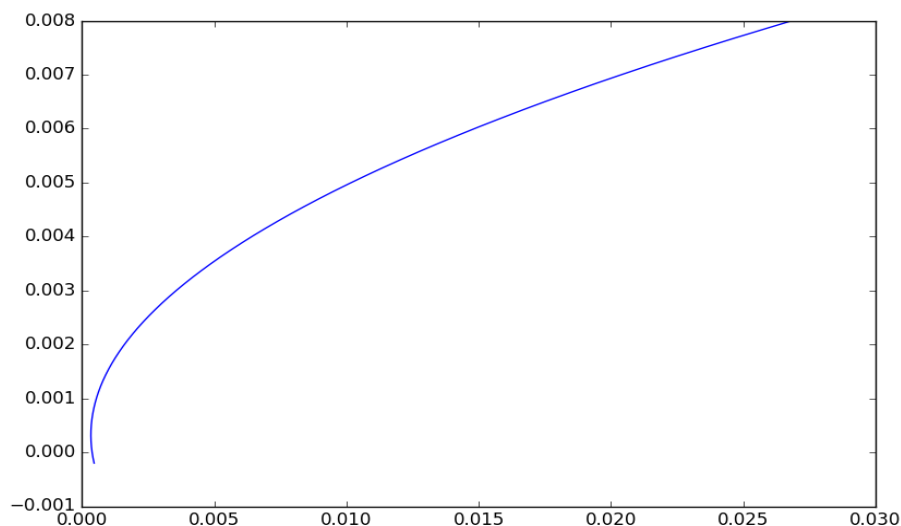


图 19.1: 题 5

```

1 In [10]: %paste
2 goal_return=0.05
3 portfolio_weight=minVar.minVar(goal_return)
4 portfolio_weight
5
6 ## — End pasted text —
7 Out[10]:
8 array([[ ' sclq' , ' hsl' , ' wld' , ' ghn' , ' zhj' ],
9        [ 9.040199555721635 , ' -44.98072150897727 , ' 9.347995192679459 ,
10         ' 24.975718390028632 , ' 2.6168083705475422 ]],
11        dtype='<U32 ')

```

6. 选择矩阵:

	A	B	C	D
1	0	0	0	0
0	-1	0	1	
1	0	-1	0	

看法向量: (0.06, 0.02, 0.03)

信心向量: (0.85, 0.90, 0.75)

8. 计算步骤与文中一样。

```

1 In [1]: %paste
2 import numpy as np
3 import math
4 import matplotlib.pyplot as plt
5 import pandas as pd

```

```

6 import pandas_datareader.data as web
7 import datetime as dt
8 import ffn
9 from scipy import linalg
10
11 ## — End pasted text —
12
13 In [2]: %paste
14 def blacklitterman(returns, tau, P, Q):
15     mu=returns.mean()
16     sigma=returns.cov()
17     pi1=mu
18     ts = tau * sigma
19     Omega = np.dot(np.dot(P, ts), P.T) * np.eye(Q.shape[0])
20     middle = linalg.inv(np.dot(np.dot(P, ts), P.T) + Omega)
21     er = np.expand_dims(pi1, axis=0).T + np.dot(np.dot(np.dot(ts, P.T), middle),
22         (Q - np.expand_dims(np.dot(P, pi1.T), axis=1)))
23     posteriorSigma = sigma + ts - np.dot(ts.dot(P.T).dot(middle).dot(P), ts)
24     return [er, posteriorSigma]
25
26 def blminVar(blres, goalRet):
27     covs=np.array(blres[1])
28     means=np.array(blres[0])
29     L1=np.append(np.append((covs.swapaxes(0,1)), [means.flatten()], 0),
30         [np.ones(len(means)), 0].swapaxes(0,1)
31     L2=list(np.ones(len(means)))
32     L2.extend([0,0])
33     L3=list(means)
34     L3.extend([0,0])
35     L4=np.array([L2, L3])
36     L=np.append(L1, L4, 0)
37     results=linalg.solve(L, np.append(np.zeros(len(means)), [1, goalRet], 0))
38     return(pd.DataFrame(results[:-2],
39         index=blres[1].columns, columns=[ 'p_weight' ]))
40 ## — End pasted text —
41
42 In [3]: %paste
43 pick1 = np.array([1,0,0,0])
44 pick2 = np.array([0,-1,0,1])
45 pick3 = np.array([1,0,-1,0])

```

```

46 P = np.array([pick1,pick2,pick3])
47 Q=np.array([[0.06],[0.02],[0.03]])
48
49 ## — End pasted text —
50
51 In [4]: %paste
52 lcxx = web.DataReader(' 000977.SZ ', yahoo ,
53                       dt.datetime(2012,1,2),dt.datetime(2014,12,31)).Close
54 lcxxRet=((lcxx-lcxx.shift(1))/lcxx.shift(1))
55 lcxxRet.name=' lcxx'
56
57 pfyh = web.DataReader(' 600000.SS ', yahoo ,
58                       dt.datetime(2012,1,2),dt.datetime(2014,12,31)).Close
59 pfyhRet=((pfyh-pfyh.shift(1))/pfyh.shift(1))
60 pfyhRet.name=' pfyh'
61
62 htzq = web.DataReader(' 601688.SS ', yahoo ,
63                       dt.datetime(2012,1,2),dt.datetime(2014,12,31)).Close
64 htzqRet=((htzq-htzq.shift(1))/htzq.shift(1))
65 htzqRet.name=' htzq'
66
67 zqb = web.DataReader(' 300052.SZ ', yahoo ,
68                      dt.datetime(2012,1,2),dt.datetime(2014,12,31)).Close
69 zqbRet=((zqb-zqb.shift(1))/zqb.shift(1))
70 zqbRet.name=' zqb'
71
72 ## — End pasted text —
73
74 In [5]: %paste
75 Stocks=pd.concat([lcxxRet,pfyhRet,htzqRet,zqbRet],1)
76 Stocks=Stocks.dropna()
77
78 ## — End pasted text —
79
80 In [6]: %paste
81 res=blacklitterman(Stocks,0.1, P, Q)
82 blminVar(res,0.75/252)
83
84 ## — End pasted text —
85 Out[6]:

```

```

86      p_weight
87 lcxx -0.015149
88 pfyh 0.810286
89 htzq 0.082629
90 zqb 0.122235

```

9. 假设投资组合为 $w = \sum_{i=1}^5 w_i x_i$, 那么我们需要求解的方差最小化问题则为:

$$\min Var(w)$$

subject to

$$\begin{cases} \sum_{i=1}^5 w_i = 1 \\ \sum_{i=1}^5 w_i \bar{R}_i = 0.05 \end{cases}$$

接着, 我们可以建立如下的拉格朗日函数:

$$L(w_1, w_2, w_3, w_4, w_5, \lambda_1, \lambda_2) = Var(w) - \lambda_1 \left(\sum_{i=1}^5 w_i - 1 \right) - \lambda_2 \left(\sum_{i=1}^5 w_i \bar{R}_i - 0.05 \right)$$

因为 $Var(w) = \sum_{i=1}^5 w_i^2 Var(R_i) + \sum_{i=1}^5 \sum_{j=1, j \neq i}^5 w_i w_j Cov(R_i, R_j)$, 所以我们只需求解下列方程组即可得到方差最小化的组合权重。¹

$$\begin{cases} 2Var(R_i)w_i + 2\sum_{j=1, j \neq i}^5 w_j Cov(R_i, R_j) - \lambda_1 - \lambda_2 \bar{R}_i = 0 & i = 1, 2, 3, 4, 5 \\ \sum_{i=1}^5 w_i = 1 \\ \sum_{i=1}^5 w_i \bar{R}_i = 0.05 \end{cases}$$

¹严格地说, 该方程组的解仅是可能的极值点。但在该例中, 方程组的解不仅是极小值, 而且是最小值。

第二十章 资本资产定价模型

1.

(a) $R = R_f + (R_m - R_f) = 10\%$

(b) $R = R_f + 1.5(R_m - R_f) = 13\%$

(c) $R = R_f + 0.7(R_m - R_f) = 8.2\%$

2. A、B、C 都被错误地定价了。

3.

(a) 无影响

(b) 该事件提高了公司的经营风险，从而提高公司的 β 值。

(c) 该事件降低了公司的财务杠杆，从而降低了公司的 β 值。

4.

(a) 该公司的系统风险与市场组合风险变动的幅度和方向一致

(b) 该公司的系统风险与市场组合的风险无关

5. 风险厌恶的投资者应该选择 β 值为 0.7 的股票，因为 β 值越大表明公司的系统风险相对于市场组合风险就越大。

6. $\beta = \frac{Cov(R_A, R_m)}{Var(R_m)} = \frac{0.96}{0.66} = 1.45$

7. A 公司的 α 值为: $9\% - (4\% + 1.2 \times 6\%) = -2.2\%$

B 公司的 α 值为: $12\% - (4\% + 1.3 \times 6\%) = 0.2\%$

所以，投资者应该选取 B 公司的股票。

8.

```
1 In [1]: import pandas as pd
2     ...:
3     ...: import pandas_datareader.data as web
4     ...:
5     ...: import datetime as dt
6     ...:
7     ...: import statsmodels.formula.api as smf
8
9 In [2]: rf=1.036**(1/360)-1
```

```

10
11 In [3]: nyh = web.DataReader(' 601288.SS ', yahoo ,
12     ...:                      dt.datetime(2014,1,1),dt.datetime(2014,12,31))
13
14 In [4]: returns = (nyh.Close-nyh.Close.shift(1))/nyh.Close.shift(1)
15
16 In [5]: indexcd=pd.read_csv(' Data/Part3/003/TRD_Index.csv ',index_col= Trddt )
17
18 In [6]: mktcd=indexcd[indexcd.Indexcd==902]
19
20 In [7]: mktret=pd.Series(mktcd.Retindex.values,index=pd.to_datetime(mktcd.index))
21
22 In [8]: mktret.name= mktret
23
24 In [9]: mktret=mktret[' 2014-01-02 ': 2014 ]
25
26 In [10]: dat = pd.concat([mktret,returns],1)
27
28 In [11]: dat = dat - rf
29
30 In [12]: model = smf.ols(' Close~mktret ',data=dat).fit()
31
32 In [13]: print(model.summary())

```

```

=====
                        OLS Regression Results
=====
Dep. Variable:          Close      R-squared:                0.118
Model:                  OLS      Adj. R-squared:             0.114
Method:                 Least Squares      F-statistic:         32.41
Date:                  Tue, 17 Jan 2017      Prob (F-statistic):    3.59e-08
Time:                  20:47:27      Log-Likelihood:       661.60
No. Observations:      245      AIC:                  -1319.
Df Residuals:          243      BIC:                  -1312.
Df Model:               1
Covariance Type:       nonrobust
=====

```

	coef	std err	t	P> t	[95.0% Conf. Int.]
Intercept	0.0009	0.001	0.849	0.397	-0.001 0.003
mktret	0.5339	0.094	5.693	0.000	0.349 0.719

```

=====
Omnibus:                130.839      Durbin-Watson:         1.954
Prob(Omnibus):           0.000      Jarque-Bera (JB):      2133.880
Skew:                   1.687      Prob(JB):              0.00
Kurtosis:                17.059      Cond. No.              89.9
=====

```

所以，模型为

$$R_i - R_f = 0.0009 + 0.5339 \times (R_m - R_f) + \varepsilon$$

其中， $R_f = 1.036^{1/360} - 1$ 。

1.

(a)

```
1 In [14]: from statsmodels.api import add_constant
2
3 In [15]: rf=1.036**(1/360)-1
4
5 In [16]: lsw = web.DataReader(' 300104.SZ ', yahoo ,
6     ...:                      dt.datetime(2014,1,1),dt.datetime(2014,12,31))
7
8 In [17]: returns = (lsw.Close-lsw.Close.shift(1))/lsw.Close.shift(1)
9
10 In [18]: indexcd=pd.read_csv(' Data/Part3/003/TRD_Index.csv ',index_col= 'Trddt'
11     )
12
13 In [19]: mktcd=indexcd[indexcd.Indexcd==902]
14
15 In [20]: mktret=pd.Series(mktcd.Retindex.values ,index=pd.to_datetime(mktcd.
16     index))
17
18 In [21]: mktret.name= 'mktret'
19
20 In [22]: mktret=mktret[' 2014-01-02 ': 2014 ]
21
22 In [23]: dat = pd.concat([mktret,returns],1)
23
24 In [24]: dat = dat - rf
25
26 In [25]: model = smf.ols(' Close~mktret ',data=dat).fit()
27
28 In [26]: print(model.summary())
```

```

=====
                        OLS Regression Results
=====
Dep. Variable:          Close      R-squared:                0.079
Model:                  OLS       Adj. R-squared:           0.075
Method:                 Least Squares   F-statistic:             20.77
Date:                  Tue, 17 Jan 2017   Prob (F-statistic):       8.21e-06
Time:                  20:49:55    Log-Likelihood:           504.68
No. Observations:      245         AIC:                     -1005.
Df Residuals:          243         BIC:                     -998.4
Df Model:               1
Covariance Type:       nonrobust
=====

```

	coef	std err	t	P> t	[95.0% Conf. Int.]
Intercept	-0.0017	0.002	-0.870	0.385	-0.006 0.002
mktret	0.8108	0.178	4.557	0.000	0.460 1.161

```

=====
Omnibus:                27.033   Durbin-Watson:           1.956
Prob(Omnibus):           0.000   Jarque-Bera (JB):         57.347
Skew:                    0.545   Prob(JB):                  3.53e-13
Kurtosis:                5.105   Cond. No.:                 89.9
=====

```

(b)

```

1 In [27]: ret_2015=pd.Series(mktcd.Retindex.values,index=pd.to_datetime(mktcd.
      index))[2015-01-01:]
2
3 In [28]: ret_2015.name='mktret'
4
5 In [29]: ret_2015 = (ret_2015 - rf)
6
7 In [30]: prediction = model.predict(add_constant(ret_2015),transform=False)

```

读者可自行查询乐视的股价。总的来说，预测值与实际值的差距还是比较大的。这说明市场回报率是无法完全解释个股的回报率，我们还需要其他因素。下一章的 Fama-French 三因子模型即是一个对资本资产定价模型的拓展。

2. 由于股票众多，在此以银行产业的股票作为范例。从 `Stock.acddb` 文件中选择银行产业的股票分别建立 CAPM 模型，并筛选出 2013 年 `alpha` 值最大的前 3 只股票。

```

1 In [1]: %paste
2 import pypyodbc
3 import pandas as pd
4 import datetime as dt
5 import numpy as np
6 import statsmodels.formula.api as smf
7 from statsmodels.api import add_constant
8
9 ## — End pasted text —

```

(a) 获取银行产业股票代码。

```

1 In [2]: %paste

```

```

2 for line in open( Data/Part3/003/industryCodes.txt ):
3     listData=line.split( ';' )
4     fileName=listData[0]
5     listData[-1]=listData[-1][:6]
6     if fileName!=' 银行 ':
7         break
8
9 ## — End pasted text —

```

(b) 以中证流通收益率数据作为市场收益率数据。

```

1 In [3]: %paste
2 indexcd=pd.read_table( Data/Part3/003/TRD_Index.txt , sep=' \t ' ,index_col='
   Trddt ' )
3 mktcd=indexcd[indexcd.Indexcd==902]
4 mktret=pd.Series( mktcd.Retindex.values , index=pd.to_datetime( mktcd.index ) )
5 mktret.name=' mktret '
6
7 ## — End pasted text —

```

(c) 编写函数，并对求出股票的 alpha 值。

```

1 In [4]: %paste
2 def GetStockAlpha(symbol, mktret):
3     if symbol[0]!=' 0 ':
4         i=0
5         while (symbol[i]!=' 0 '):
6             i+=1
7         symbol=symbol[i:]
8     #先提前安装Microsoft Access Database Engine
9     #可以去Microsoft下载中心下载Microsoft Access Database Engine
10    conn = pyodbc.connect(' Driver=Microsoft Access Driver (*.mdb, *.
   accdb);DBQ=%s;' %accessName )
11    cursor = conn.cursor()
12    cursor.execute(' select Stkcd, Trddt, Clsprc from stock where Stkcd=' +
   symbol+' ' )
13    data=cursor.fetchall()
14    prices=[]
15    dates=[]
16    for entry in data:
17        prices.append(entry[-1])
18        time=str(entry[1])
19        date=pd.to_datetime(time[1:])

```

```

20         dates.append(date)
21     price=pd.Series(np.array(prices),index=dates)[ 2013 ]
22     returns = (price-price.shift(1))/price.shift(1)
23     returns.name= stock
24
25     rf=1.036**(1/360)-1
26     dat = pd.concat([mktret[ 2013-01-06 : 2013 ],returns[ 2013-01-06 :
27         2013 ]],1)
28     dat = dat - rf
29     model = smf.ols(' stock~mktret ',data=dat).fit()
30     return(model.params[0])
31 ## — End pasted text —
32
33 In [4]: %paste
34 alphas={}
35 for symbol in listData[1:]:
36     try:
37         alphas[symbol]=GetStockAlpha(symbol,mktret)
38     except Exception as e:
39         print (Exception," :",e)
40
41 ## — End pasted text —

```

(d) 筛选 alpha 值前 3 大的股票代码及其 alpha 值。

```

1 In [5]: %paste
2 selectAlpha=[(' -00 ',-100000),(' -00 ',-100000),(' -00 ',-100000)]
3 for symbol in alphas.keys():
4     if alphas[symbol]>selectAlpha[0][1]:
5         selectAlpha[0]=(symbol,alphas[symbol])
6         selectAlpha.sort(key=lambda d:d[1])
7 selectAlpha
8
9 ## — End pasted text —
10 Out[5]:
11 [(' 601998 ', -0.00056736014634597726),
12  (' 600000 ', -0.00032213839451176914),
13  (' 600016 ', -0.0002028846807797023)]

```

第二十一章 三因子模型

1. 一个可能的猜想即是同行业的其他公司的股价收益率会影响该股票的收益率。下面我们来验证一下。

```
1 In [1]: import pandas as pd
2     ...:
3     ...: import pandas_datareader.data as web
4     ...:
5     ...: import datetime as dt
6     ...:
7     ...: import statsmodels.formula.api as smf
8     ...:
9     ...: from statsmodels.api import add_constant
10
11 In [2]: wanke = web.DataReader(' 000002.SZ ', yahoo ,
12     ...:                        dt.datetime(2015,1,1),dt.datetime(2015,12,31))
13
14 In [3]: gldc = web.DataReader(' 600185.SS ', yahoo ,
15     ...:                        dt.datetime(2015,1,1),dt.datetime(2015,12,31))
16
17 In [4]: price = pd.concat([wanke.Close,gldc.Close],1)
18
19 In [5]: price.columns=[ 'wanke' , 'gldc' ]
20
21 In [6]: ret = (price-price.shift(1)) / price.shift(1)
22
23 In [7]: model = smf.ols(' wanke~gldc ',data=ret).fit()
24
25 In [8]: print(model.summary())
```

OLS Regression Results						
=====						
Dep. Variable:	wanke		R-squared:	0.176		
Model:	OLS		Adj. R-squared:	0.173		
Method:	Least Squares		F-statistic:	55.25		
Date:	Tue, 17 Jan 2017		Prob (F-statistic):	1.56e-12		
Time:	20:58:05		Log-Likelihood:	568.10		
No. Observations:	260		AIC:	-1132.		
Df Residuals:	258		BIC:	-1125.		
Df Model:	1					
Covariance Type:	nonrobust					
=====						
	coef	std err	t	P> t	[95.0% Conf. Int.]	

Intercept	0.0024	0.002	1.398	0.163	-0.001	0.006
geli	0.2808	0.038	7.433	0.000	0.206	0.355
=====						
Omnibus:	43.977		Durbin-Watson:	1.947		
Prob(Omnibus):	0.000		Jarque-Bera (JB):	103.520		
Skew:	0.799		Prob(JB):	3.32e-23		
Kurtosis:	5.646		Cond. No.	22.3		
=====						

模型的结果显示变量 *geli* 的系数是显著的，说明格力地产的股价确实和万科的股价有一定的关系。我们的假设是正确的。

1. CAPM 模型和三因子模型的异同点：

- (a) 相同点：都用于通过建立线性模型来探究影响股票收益率的因素。
- (b) 不同点：
 - i. CAPM 模型只考虑了一个因子，市场因子；三因子模型考虑了市场风险溢价因子、市值因子和账面市值比因子这三种因子。
 - ii. CAPM 模型只考虑了市场风险，三因子模型除了市场风险以外，还考虑了不同公司本身的运营情况和在市场中的价值等因素，更加多方面考虑了股票收益率的影响因素。

2. $SMB_t = 1/3(SL_t + SM_t + SH_t) - 1/3(BL_t + BM_t + BH_t) = -0.03$

$HML_t = 1/2(SL_t + BL_t) - 1/2(SH_t + BH_t) = 0.1$

3.

- (a) 0.01
- (b) $R_i = R_f + 0.01 + 1.2 \times (2\% - 0.5\%) + 0.5 \times 2.4\% + 0.1 \times 1.8\% = 4.68\%$

```
In [9]: zyhy = pd.read_table(' Data/Part3/004/problem21.txt' ,
....:     sep=' \t' ,usecols=[ zyhy' ; Date' ],index_col=' Date' )
In [10]: zyhy.index=pd.to_datetime(zyhy.index)
In [11]: ret = (zyhy-zyhy.shift(1))/zyhy.shift(1)
In [12]: ThreeFactors = pd.read_table(' Data/Part3/004/ThreeFactors.txt' ,
....:     sep=' \t' ,index_col=' TradingDate' )
```

```

10
11 In [13]: ThreeFactors.index=pd.to_datetime(ThreeFactors.index)
12
13 In [14]: ThrFac=ThreeFactors[ 2014 ]
14
15 In [15]: ThrFac=ThrFac.iloc[:,[2,4,6]]
16
17 In [16]: dat=pd.concat([ret,ThrFac],1)
18
19 In [17]: model = smf.ols( zyhy~RiskPremium2+SMB2+HML2 ,data=dat).fit()
20
21 In [18]: print(model.summary())

```

```

=====
                        OLS Regression Results
=====
Dep. Variable:          zyhy      R-squared:            0.223
Model:                  OLS      Adj. R-squared:        0.213
Method:                 Least Squares    F-statistic:      23.01
Date:                  Mon, 23 Jan 2017    Prob (F-statistic):  3.94e-13
Time:                  15:56:16    Log-Likelihood:     598.75
No. Observations:      245    AIC:                  -1190.
Df Residuals:          241    BIC:                  -1175.
Df Model:               3
Covariance Type:       nonrobust
=====

```

	coef	std err	t	P> t	[95.0% Conf. Int.]
Intercept	0.0021	0.001	1.507	0.133	-0.001 0.005
RiskPremium2	0.9948	0.126	7.893	0.000	0.746 1.243
SMB2	-0.7332	0.223	-3.288	0.001	-1.172 -0.294
HML2	0.7493	0.211	3.559	0.000	0.335 1.164

```

=====
Omnibus:                105.816    Durbin-Watson:          1.766
Prob(Omnibus):           0.000    Jarque-Bera (JB):        440.407
Skew:                    1.763    Prob(JB):                2.33e-96
Kurtosis:                 8.542    Cond. No.:               220.
=====

```

```

4.
1 In [19]: zhongxin = pd.read_table( problem21.txt ,sep= \t ,usecols=[ zhongxin ,
2         Date ],index_col= Date )
3
4 In [20]: zhongxin.index=pd.to_datetime(zhongxin.index)
5
6 In [21]: ret = (zhongxin-zhongxin.shift(1))/zhongxin.shift(1)
7
8 In [22]: ThreeFactors = pd.read_table( ThreeFactors.txt ,sep= \t ,index_col=
9         TradingDate )
10
11 In [23]: ThreeFactors.index=pd.to_datetime(ThreeFactors.index)
12
13 In [24]: ThrFac=ThreeFactors[ 2014 ]
14
15

```

```

13 In [25]: ThrFac=ThrFac.iloc[:,[2,4,6]]
14
15 In [26]: dat=pd.concat([ret,ThrFac],1)

```

(a)

```

1 In [27]: model = smf.ols(' zhongxin~RiskPremium2 ',data=dat).fit()
2
3 In [28]: print(model.summary())

```

```

=====
                    OLS Regression Results
=====
Dep. Variable:          zhongxin      R-squared:                0.207
Model:                  OLS          Adj. R-squared:            0.204
Method:                 Least Squares   F-statistic:              63.51
Date:                  Mon, 23 Jan 2017   Prob (F-statistic):       6.19e-14
Time:                  15:59:44         Log-Likelihood:           557.71
No. Observations:      245             AIC:                     -1111.
Df Residuals:          243             BIC:                     -1104.
Df Model:              1
Covariance Type:       nonrobust
=====

```

	coef	std err	t	P> t	[95.0% Conf. Int.]
Intercept	0.0015	0.002	0.950	0.343	-0.002 0.005
RiskPremium2	1.1710	0.147	7.970	0.000	0.882 1.460

```

=====
Omnibus:                 89.839   Durbin-Watson:           1.716
Prob(Omnibus):            0.000   Jarque-Bera (JB):        370.027
Skew:                     1.460   Prob(JB):                4.46e-81
Kurtosis:                 8.265   Cond. No.                 92.2
=====

```

(b)

```

1 In [29]: model2 = smf.ols(' zhongxin~RiskPremium2+SMB2+HML2 ',data=dat).fit()
2
3 In [30]: print(model2.summary())

```

```

=====
                    OLS Regression Results
=====
Dep. Variable:          zhongxin      R-squared:                0.377
Model:                  OLS          Adj. R-squared:            0.369
Method:                 Least Squares   F-statistic:              48.59
Date:                  Tue, 17 Jan 2017   Prob (F-statistic):       1.35e-24
Time:                  21:05:23         Log-Likelihood:           587.21
No. Observations:      245             AIC:                     -1166.
Df Residuals:          241             BIC:                     -1152.
Df Model:              3
Covariance Type:       nonrobust
=====

```

	coef	std err	t	P> t	[95.0% Conf. Int.]
Intercept	0.0011	0.001	0.786	0.433	-0.002 0.004
RiskPremium2	1.1763	0.132	8.903	0.000	0.916 1.437
SMB2	-0.0076	0.234	-0.032	0.974	-0.468 0.453
HML2	0.8663	0.221	3.926	0.000	0.432 1.301

```

=====
Omnibus:                 84.500   Durbin-Watson:           1.780
Prob(Omnibus):            0.000   Jarque-Bera (JB):        388.225
Skew:                     1.318   Prob(JB):                4.99e-85
Kurtosis:                 8.575   Cond. No.                 220.
=====

```

(c)

```

1 In [31]: ThrFac=ThreeFactors[ 2015-01 ]
2

```

```

3 In [32]: preCAPM = model.predict(add_constant(ThrFac.RiskPremium2), transform=
         False)
4
5 In [33]: preFactors = model2.predict(add_constant(ThrFac[[' RiskPremium2 ',
         SMB2 ', HML2 ]]),
6         ..., transform = False))

```

```

1 In [1]: %paste
2 import pandas as pd
3
4 codes = pd.read_csv(' Data/Part3/004/codes.csv ', header=None, dtype=**str)
5
6 ThreeFactors = pd.read_table(' Data/Part3/004/ThreeFactors.txt ', sep=' \t ',
         index_col=' TradingDate ')
7
8 ThreeFactors.index=pd.to_datetime(ThreeFactors.index)
9
10 ThrFac=ThreeFactors[' 2014 ']
11
12 ThrFac=ThrFac.iloc[:, [2, 4, 6]]
13
14 def create_func(model):
15     def cal_alpha(code, model_name=model):
16         price = web.DataReader(code, ' yahoo ',
17                                 dt.datetime(2014, 1, 1), dt.datetime(2014, 12, 31)).Close
18
19         ret = (price-price.shift(1))/price.shift(1)
20
21         ret.name = ' ret '
22
23         dat=pd.concat([ret, ThrFac], 1)
24
25         if model_name!=' CAPM ':
26             model = smf.ols(' ret~RiskPremium2 ', data=dat).fit()
27         elif model_name!=' factors ':
28             model = smf.ols(' ret~RiskPremium2+SMB2+HML2 ', data=dat).fit()
29
30         return(model.params[0])
31     return(cal_alpha)
32

```

33 ## — End pasted text —

6. (a)

```
1 In [14]: alpha_CAPM=list(map(create_func(' CAPM '),codes[0].values))
2
3 In [15]: alpha_CAPM2=pd.Series(alpha_CAPM).sort(ascending=False,inplace=False)
4
5 In [16]: alpha_CAPM2[:3]
6 Out[16]:
7 12    0.001531
8  8    0.000900
9  9    0.000768
10 dtype: float64
```

```
1 In [17]: alpha_factors=list(map(create_func(' factors' ),codes[0]))
2
3 In [18]: alpha_factors2 = pd.Series(alpha_factors).sort(ascending=False,
4               inplace=False)
5
6 In [19]: alpha_factors2[:3]
7 Out[19]:
8 12    0.001128
9  8    0.000532
10 13    0.000476
11 dtype: float64
```

(b) 两者选股的结果是较为相似的，仅有一个不一样。

第四部分

时间序列

第二十二章 简介

1. 略

2.

```
1 In [1]: import pandas as pd
2         ...: import matplotlib.pyplot as plt
3         ...: Yen=pd.read_csv(' Data/Part4/001/Yen.csv' ,index_col= date )
4
5 In [2]: Yen.index=pd.to_datetime(Yen.index ,format= '%Y-%m-%d' )
```

3. 对于 xts 对象, plot() 一次只能绘制一个变量。因此, 我们选择变量 *s* 为例。可以看到, 图 22.1 中存在着明显的下降的趋势。

```
1 In [3]: Yen.s.plot()
```

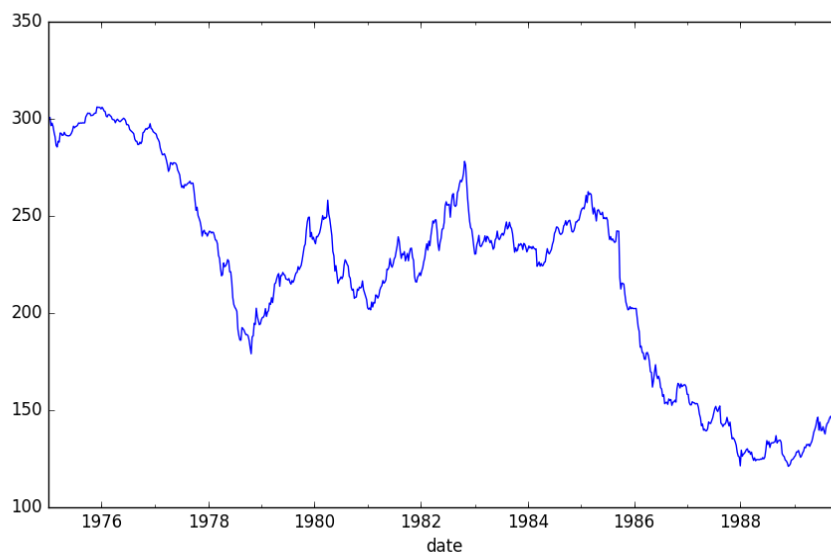


图 22.1: 题 3

4.

```
1 In [4]: pfyh=pd.read_csv(' Data/Part4/001/pfyh.csv' ,index_col= Date )
2
3 In [5]: pfyh.index=pd.to_datetime(pfyh.index ,format= '%Y-%m-%d' )
4
5 In [6]: returns=(pfyh.Close-pfyh.Close.shift(1))/pfyh.Close.shift(1)
```

```
6  
7 In [7]: returns=returns[1:]  
8  
9 In [8]: returns.plot()  
10 Out[8]: <matplotlib.axes._subplots.AxesSubplot at 0x7f8c47c6b358>
```

浦发银行的收益率围绕着 0 上下波动。

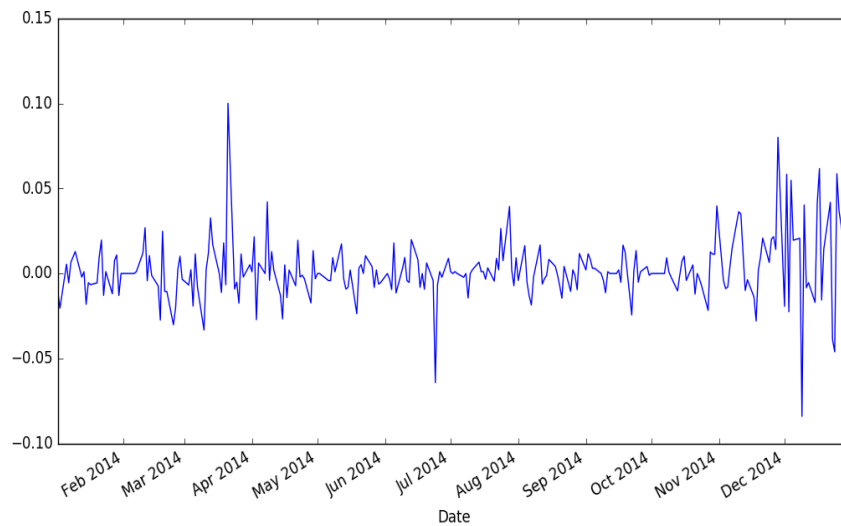


图 22.2: 题 4

第二十三章 基本性质

1.

(a) 是。当 $k = 1$ 时, $r_k = \text{corr}(X_t, X_{t-k}) = \text{corr}(e_t - e_{t-1}, e_{t-1} - e_{t-2}) = -\frac{1}{2}$;

当 $k > 1$ 时, $r_k = \text{corr}(X_t, X_{t-k}) = \text{corr}(e_t - e_{t-1}, e_{t-k} - e_{t-k-1}) = 0$

(b) 是, 其自相关系数皆为 0。

$$r_k = \text{corr}(X_t, X_{t-k}) = \text{corr}(3t, 3(t-k)) = 0$$

(c) 是, 其自相关系数皆为 0。

$$r_k = \text{corr}(X_t, X_{t-k}) = \text{corr}\left((-1)^t e_t, (-1)^{t-k} e_{t-k}\right) = (-1)^{2t-k} \text{corr}(e_t, e_{t-k}) = 0$$

2. 因为 $X_t = Y$, 所以在任意时点, X_t 的值都与 Y 的值一样。因此, 对于任意正整数 n , 任取 $t_1, t_2, \dots, t_n \in T$, 对任意整数 τ 有:

$$F_X(x_{t_1}, x_{t_2}, \dots, x_{t_n}) = F_Y(x_{t_1}) = F_X(x_{t_1+\tau}, x_{t_2+\tau}, \dots, x_{t_n+\tau})$$

所以 X_t 是严平稳的。又因为 Y 存在均值和方差, 所以 X_t 又一定是弱平稳的。

3.

$$r_k = \begin{cases} \frac{\text{Cov}(X_t, X_{t-1})}{\sqrt{\text{Var}(X_t)\text{Var}(X_{t-1})}} = \frac{\text{Cov}(e_t - 0.4e_{t-1} + 0.3e_{t-2}, e_{t-1} - 0.4e_{t-2} + 0.3e_{t-3})}{1.25\sigma^2} = -0.416, & k = 1 \\ \frac{\text{Cov}(X_t, X_{t-2})}{\sqrt{\text{Var}(X_t)\text{Var}(X_{t-2})}} = \frac{\text{Cov}(e_t - 0.4e_{t-1} + 0.3e_{t-2}, e_{t-2} - 0.4e_{t-3} + 0.3e_{t-4})}{1.25\sigma^2} = 0.24, & k = 2 \\ \frac{\text{Cov}(X_t, X_{t-k})}{\sqrt{\text{Var}(X_t)\text{Var}(X_{t-k})}} = \frac{\text{Cov}(e_t - 0.4e_{t-1} + 0.3e_{t-2}, e_{t-k} - 0.4e_{t-k-1} + 0.3e_{t-k-2})}{1.25\sigma^2} = 0, & k > 2 \end{cases}$$

显然, X_t 的自相关系数与 t 无关, 所以 X_t 是弱平稳的。

4. $\{X_t\}$ 的滞后算子多项式方程为

$$1 - x = 0$$

该方程唯一的解为 2, 大于 1。因此, $\{X_t\}$ 是一个平稳序列。其自相关系数为:

$$r_k = 0.5^k$$

5.

```
1 In [1]: import pandas as pd
2
3 In [2]: import matplotlib.pyplot as plt
```

```
4
5 In [3]: CRSPday=pd.read_csv(' Data/Part4/002/CRSPday.csv ')
6
7 In [4]: ibm=CRSPday.ibm
```

(a)

```
1 In [5]: ibm.plot()
2 Out[5]: <matplotlib.axes._subplots.AxesSubplot at 0x7f09015cb240>
```

```
1 In [6]: from statsmodels.graphics.tsaplots import *
2 In [7]: plot_acf(ibm,lags=20)
3 Out[7]: <matplotlib.figure.Figure at 0x7fe634701cf8>
```

从 ACF 图中，我们可以看到，所有的自相关系数都不显著，所以 IBM 的日回报率很有可能是白噪声。

(b)

```
1 In [8]: from statsmodels.tsa import stattools
2 In [9]: LjungBox=stattools.q_stat(stattools.acf(ibm)[1:13],len(ibm))
3 In [10]: LjungBox[1][-1]
4 Out[10]: 0.47767342041175087
```

Ljung-Box 检验的 p 值为 0.4777，大于 0.05，同样表明 IBM 的日回报率是白噪声。

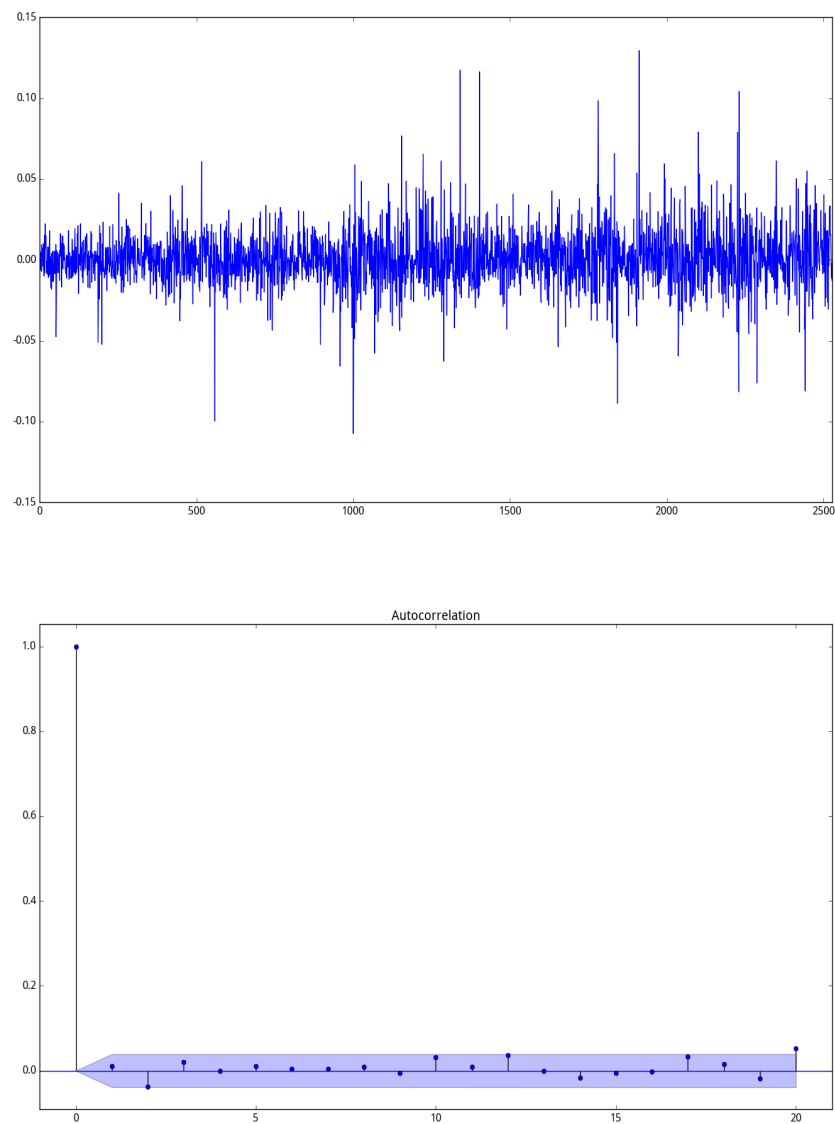


图 23.1: 题 5

6.

```
1 In [11]: ge=CRSPday.iloc[:,3]
```

(a)

```
1 In [12]: ge.plot()
```

```
2 Out[12]: <matplotlib.axes._subplots.AxesSubplot at 0x7f08e54a9da0>
```

```
1 In [13]: plot_acf(ge, lags=20)
```

(b)

```
1 In [14]: LjungBox=stattools.q_stat(stattools.acf(ge)[1:2], len(ge))
```

```
2 In [15]: LjungBox[1][-1]
```

```
3 Out[15]: 0.63415629590274136
```

```
1 In [16]: LjungBox=stattools.q_stat(stattools.acf(ge)[1:9], len(ge))
```

```
2 In [17]: LjungBox[1][-1]
3 Out[17]: 0.038343684289074531
```

- (d) 当我们使用 2 阶的自相关系数的时候, Ljung-Box 检验的 p 值为 0.6342。而当我们使用 9 阶的自相关系数的时候, Ljung-Box 检验的 p 值为 0.0383。两者的结论完全相反, 这表明选取的滞后项过少会导致检验的 p 值偏高。因此, 在实际使用 Ljung-Box 检验的时候, 我们一般会将 6 – 15 阶的自相关系数。

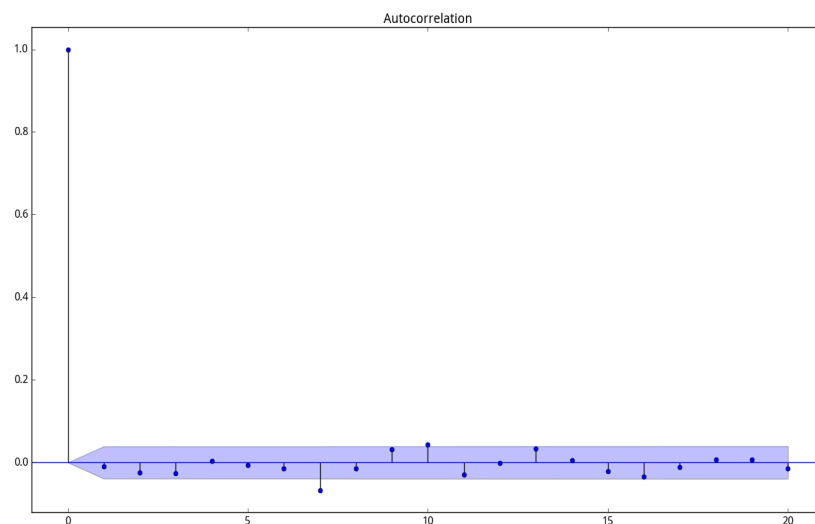


图 23.2: 题 6

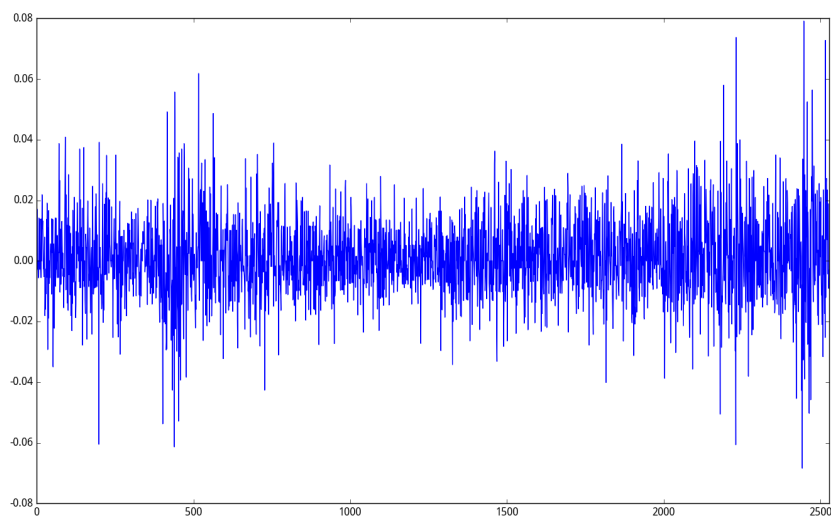


图 23.3: 题 6

```
7.
1 In [18]: SP500=pd.read_csv(' Data/Part4/002/SP500.csv ')
2
3 In [19]: r500=SP500.r500
```

(a)

```
1 In [20]: r500.plot()
2 Out[20]: <matplotlib.axes._subplots.AxesSubplot at 0x7f1593321128>
```

标普 500 指数的日回报率大致围绕着 0 上下波动，虽然中间有几个点波动过大，但总体来说还是较为平稳的。

(b)

```
1 In [21]: plot_acf(r500,lags=20)
2 Out[21]: <matplotlib.figure.Figure at 0x7fe62ee44780>
3
4 In [22]: plot_pacf(r500,lags=20)
5 Out[22]: <matplotlib.figure.Figure at 0x7fe62ee54cc0>
```

```
1 In [23]: from arch.unitroot import ADF
2
3 In [24]: adf=ADF(r500,lags=3)
4
5 In [25]: print(adf.summary().as_text())
6     Augmented Dickey-Fuller Results
7     =====
8 Test Statistic                -28.096
9 P-value                      0.000
10 Lags                        3
11
12
13 Trend: Constant
14 Critical Values: -3.43 (1%), -2.86 (5%), -2.57 (10%)
15 Null Hypothesis: The process contains a unit root.
16 Alternative Hypothesis: The process is weakly stationary.
```

ADF 检验的统计量的值为 -28.096 ，远小于 5% 的临界值 -2.86 。因此，标普 500 指数的日回报率是平稳序列。我们在 (a) 中的判断是正确的。

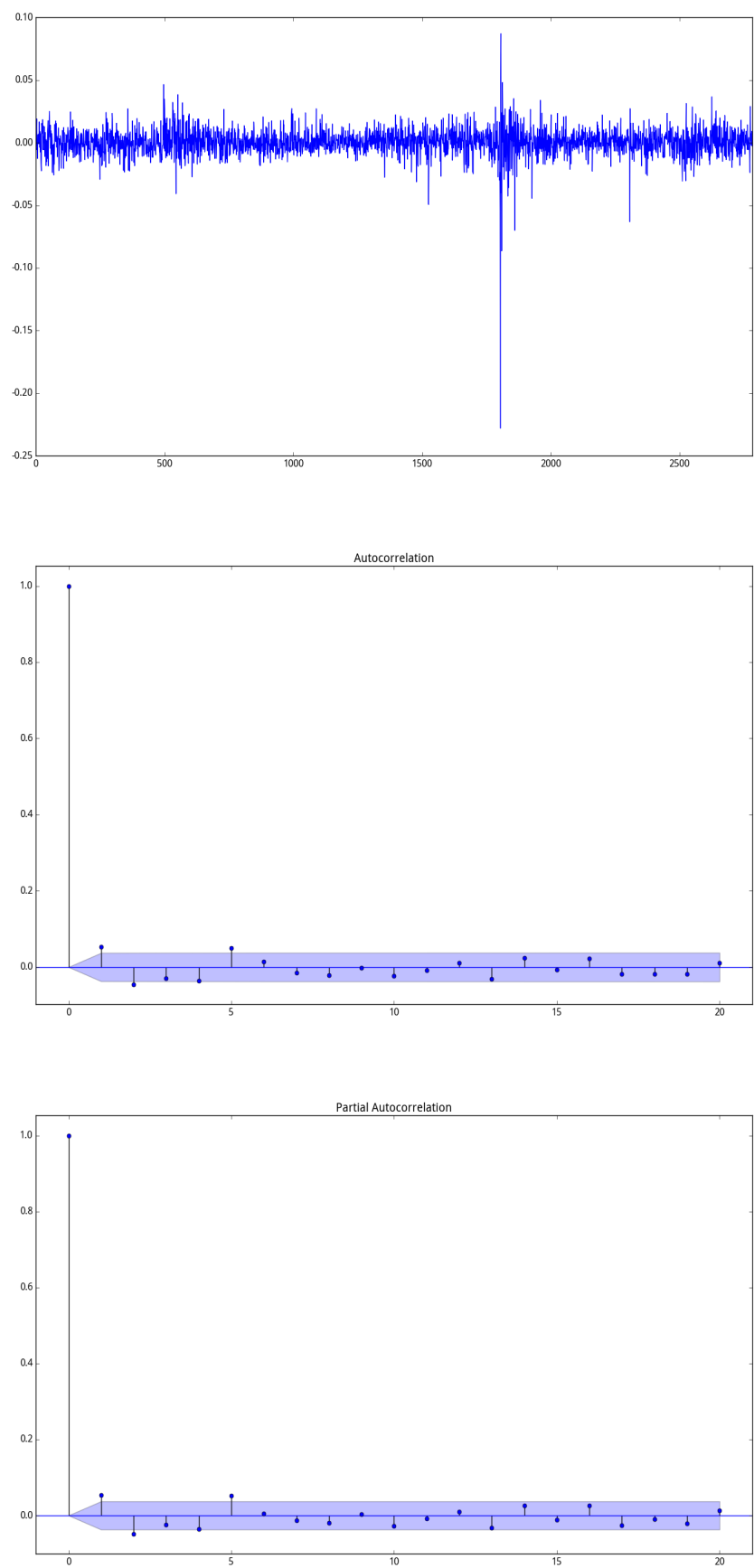


图 23.4: 题 7

第二十四章 预测

(d)

- (a) MA 模型, $q = 2$
- (b) MA 模型, $q = 4$
- (c) ARMA 模型, $p = 2, q = 3$
- (d) AR 模型, $p = 2$

2. 假设 X_t 一个 MA(q) 的时间序列, 那么 $X_t = \sum_{i=0}^q \theta_i e_{t-i}$ 。当 $k > q$ 时:

$$\rho_k = \frac{Cov(X_t, X_{t-k})}{Var(X_t)} = \frac{Cov(\sum_{i=0}^q \theta_i e_{t-i}, \sum_{i=0}^q \theta_i e_{t-k-i})}{Var(X_t)}$$

因为当 $i \neq j$ 时, $Cov(e_i, e_j) = 0$ 。所以, $\rho_k = 0$ 。

3. 第 5 期的预测值为 7.45875。

```
4.
1 In [1]: import statsmodels.tsa.arima_process as sm
2
3 In [2]: from statsmodels.graphics.tsaplots import *
4
5 In [3]: arma=sm.ArmaProcess([-1, -0.6],[1])
6
7 In [4]: sample=arma.generate_sample(200)
8
9 In [5]: plot_acf(sample, lags=20)
10 Out[5]: <matplotlib.figure.Figure at 0x7ffb9b7bfc88>
11
12 In [6]: plot_pacf(sample, lags=20)
13 Out[6]: <matplotlib.figure.Figure at 0x7ffb78ac8240>
```

ACF 图呈现出指数递减的趋势, 而 PACF 图则在二阶滞后处出现截断现象。因此, 根据绘制出来的 ACF 图与 PACF 图, 我们应该选择 AR(1) 模型。该模型符合真实情况。

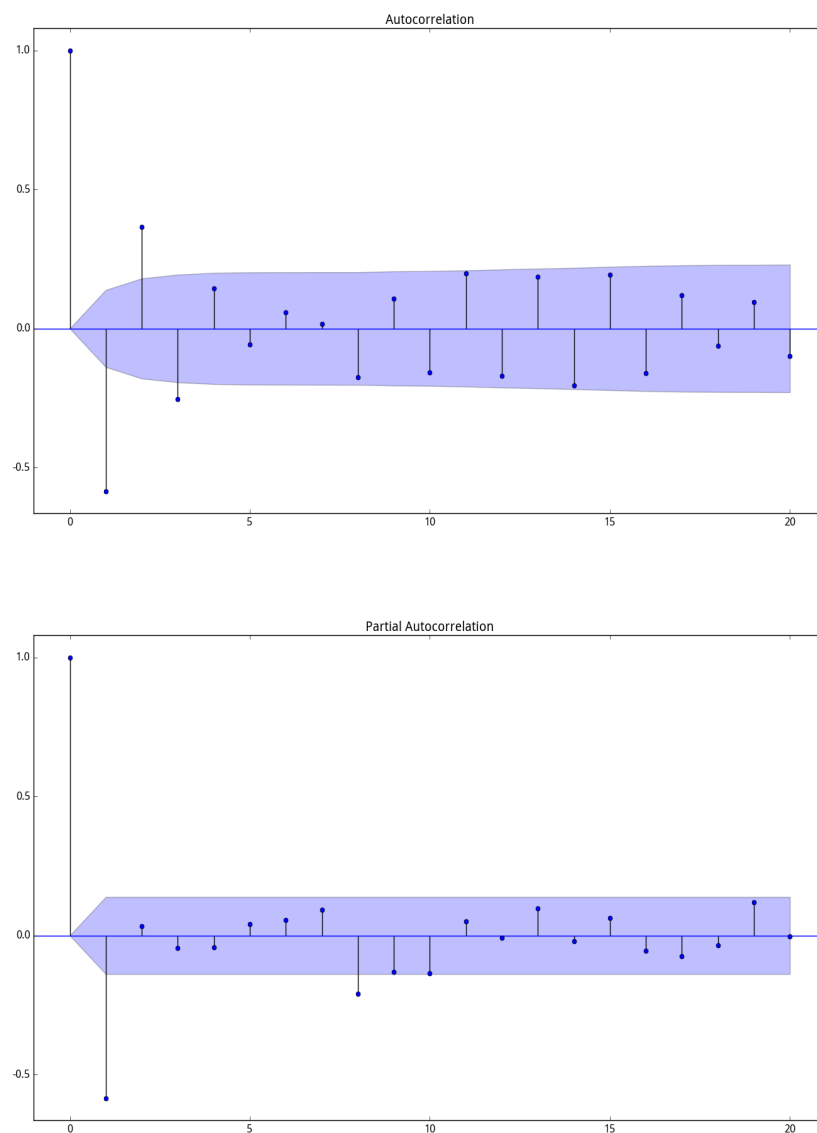


图 24.1: 题 4

```

5. In [7]: arma=sm.ArmaProcess([-1],[1,0.4])
6. In [8]: sample=arma.generate_sample(200)
7. In [9]: plot_acf(sample,lags=20)
8. Out[9]: <matplotlib.figure.Figure at 0x7ffb78ae1470>
9. In [10]: plot_pacf(sample,lags=20)
10. Out[10]: <matplotlib.figure.Figure at 0x7ffb749feac8>

```

在 ACF 图中，一阶自相关系数与二阶自相关系数都是显著的。而在 PACF 图中，一阶偏自相关系数是显著的，之后的偏自相关系数都不显著。因此，根据 ACF 图和 PACF 图，我们可能会选择 MA(2) 或是 AR(1)。这时，选择的模型就不再符合真实情况了。这是因为模拟的期数过少，样本统计量存在较大的偏差。在实务中，我们也要

注意这种偏差，多尝试几个模型，而不是一味地依照 ACF 与 PACF。比如说，在本题中，我们就可以再多尝试 ARMA(1,1) 与 MA(1) 这两个模型。

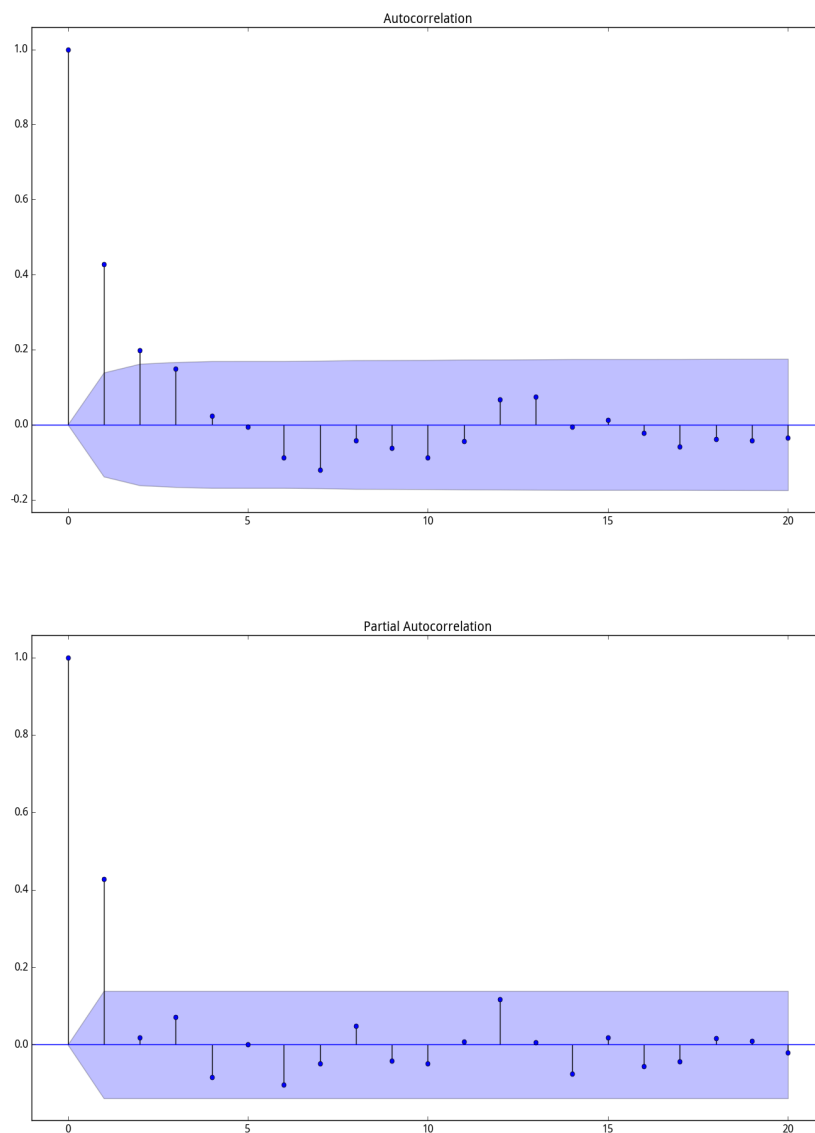


图 24.2: 题 5

```
6. In [1]: import statsmodels.tsa.arima_process as sm
1 In [2]: from statsmodels.graphics.tsaplots import *
2 In [3]: import numpy as np
3 In [4]: import pandas as pd
4 In [5]: numbers=np.random.normal(size=100)
```

10

```
11 In [6]: numbers=pd.Series(numbers)
```

(a)

```
1 In [7]: numbers.plot()
2 Out[7]: <matplotlib.axes._subplots.AxesSubplot at 0x7f24c58a2358>
3
4 In [8]: plot_acf(numbers,lags=20)
5 Out[8]: <matplotlib.figure.Figure at 0x7f24c528b828>
```

很明显, *numbers* 围绕着 0 上下随机波动, 而且其 ACF 都不显著。因此, 我们可以认为 *numbers* 是白噪声。

(b)

```
1 In [9]: from statsmodels.tsa import stattools
2
3 In [10]: stattools.arma_order_select_ic(numbers.values,max_ma=4)
4 Out[10]:
5 { bic :           0           1           2           3           4
6 0  301.560380  305.113771  309.241432  313.831146  318.394066
7 1  305.268515  309.372041  313.826739  318.431879  322.992102
8 2  309.240906  313.835172  318.902116  323.034685  324.016683
9 3  313.840677  318.439734           NaN  327.099384           NaN
10 4  318.346121  323.058940  324.742326  331.636021  334.205935,
11 'bic_min_order' : (0, 0)}
```

`arma_order_select_ic()` 函数得到的结果为 ARIMA(0,0,0), 也就是说, 我们只能使用均值来预测 *numbers* 的值。可见, 对于白噪声, 我们无法进行有意义的建模。

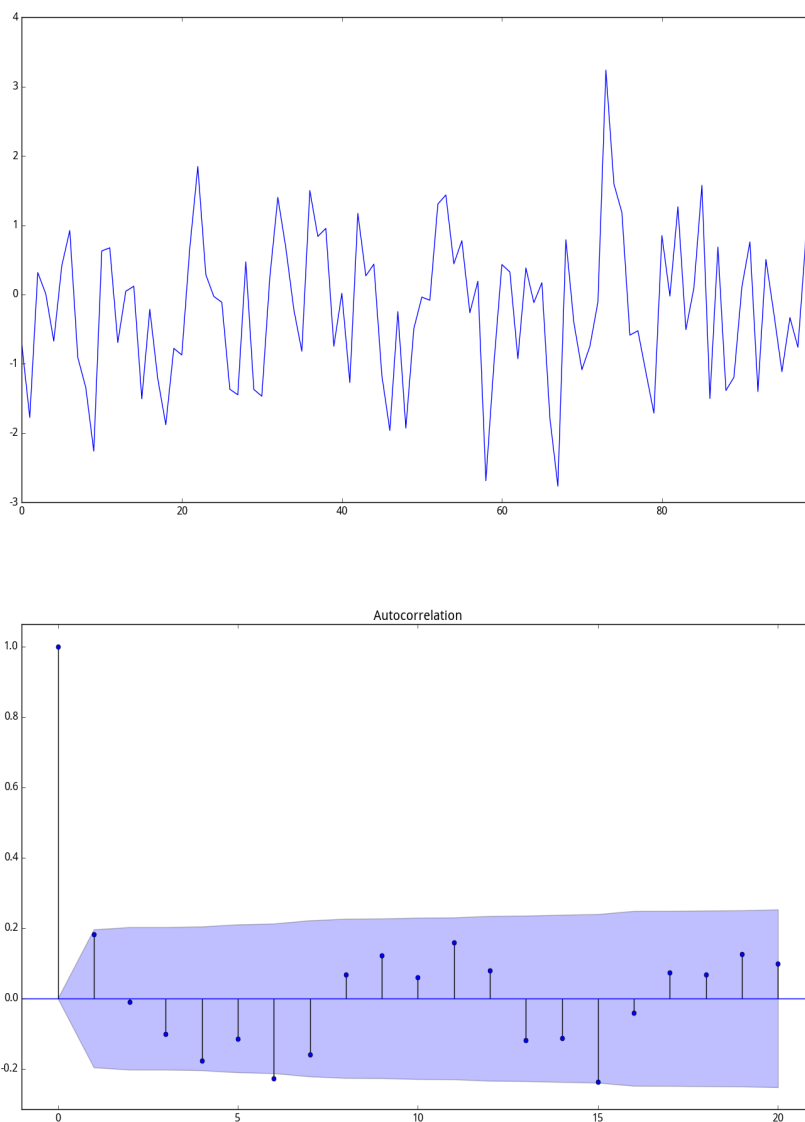


图 24.3: 题 6

```

7. In [1]: zgsy=pd.read_csv( Data/Part4/003/zgsy.csv )
    In [2]: clprice=zgsy.iloc[:,4]

```

```

(a) In [3]: clprice.plot()
      Out[3]: <matplotlib.axes._subplots.AxesSubplot at 0x7f24e3de8160>
      In [4]: plot_acf(clprice ,lags=20)
      Out[4]: <matplotlib.figure.Figure at 0x7f24e3cb3f98>

```

clprice 的时间序列图呈现出明显的下降趋势，而且其 **ACF** 也是缓慢递减。因此，*clprice* 不平稳。

```

(b) In [5]: from arch.unitroot import ADF

```

```

2     ...: adf=ADF(clprice, lags=6)
3
4 In [6]: print(adf.summary().as_text())
5     Augmented Dickey-Fuller Results
6     =====
7 Test Statistic          -1.185
8 P-value                0.680
9 Lags                   6
10
11
12 Trend: Constant
13 Critical Values: -3.46 (1%), -2.87 (5%), -2.57 (10%)
14 Null Hypothesis: The process contains a unit root.
15 Alternative Hypothesis: The process is weakly stationary.

```

ADF 检验的统计量为 -1.185 ，大于 5% 的临界值。因此，我们不能拒绝原假设，*clprice* 不平稳。

(c)

```

1 In [7]: logReturn=pd.Series((np.log(clprice))).diff().dropna()
2
3 In [8]: logReturn.plot()
4 Out[8]: <matplotlib.axes._subplots.AxesSubplot at 0x7f24e3cd1cf8>

```

可以看到，*logReturn* 围绕着 0 上下波动，并没有呈现出明显的趋势。

(d)

```

1 In [9]: adf=ADF(logReturn, lags=6)
2
3 In [10]: print(adf.summary().as_text())
4     Augmented Dickey-Fuller Results
5     =====
6 Test Statistic          -6.402
7 P-value                0.000
8 Lags                   6
9
10
11 Trend: Constant
12 Critical Values: -3.46 (1%), -2.87 (5%), -2.57 (10%)
13 Null Hypothesis: The process contains a unit root.
14 Alternative Hypothesis: The process is weakly stationary.

```

ADF 检验的结果表明我们可以拒绝原假设，因此 *logReturn* 是平稳的。

(e)

```

1 In [11]: plot_acf(logReturn, lags=20)
2 Out[11]: <matplotlib.figure.Figure at 0x7f24e3c19748>

```

```

3
4 In [12]: plot_pacf(logReturn, lags=20)
5 Out[12]: <matplotlib.figure.Figure at 0x7f24e3b7c828>

```

在 ACF 图中，自相关系数从二阶开始变得不显著。在 PACF 图中，偏自相关系数同样从二阶开始变得不显著。因此，我们可以有两个选择——MA(1) 和 AR(1)。

(f)

```

1 In [13]: from statsmodels.tsa import arima_model
2
3 In [14]: model1=arima_model.ARIMA(logReturn.values, order=(0,0,1)).fit()

```

```

RUNNING THE L-BFGS-B CODE
***
Machine precision = 2.220D-16
N = 2 M = 12
This problem is unconstrained.

At X0 0 variables are exactly at the bounds

At iterate 0 f= -3.28295D+00 |proj g|= 4.68026D-04

At iterate 5 f= -3.28295D+00 |proj g|= 9.19265D-06
***
Tit = total number of iterations
Tnf = total number of function evaluations
Tnint = total number of segments explored during Cauchy searches
Skip = number of BFGS updates skipped
Nact = number of active bounds at final generalized Cauchy point
Projg = norm of the final projected gradient
F = final function value
***
N Tit Tnf Tnint Skip Nact Projg F
2 6 16 1 0 0 3.864D-06 -3.283D+00
F = -3.2829527560393754

CONVERGENCE: REL_REDUCTION_OF_F_<=_FACTR*EPSMCH

Cauchy time 0.000E+00 seconds.
Subspace minimization time 0.000E+00 seconds.
Line search time 0.000E+00 seconds.

Total User time 0.000E+00 seconds.

```

```

1 In [15]: model1.summary()
2 Out[5]:
3 <class 'statsmodels.iolib.summary.Summary' >

```

```

"""
                                ARMA Model Results
=====
Dep. Variable: y No. Observations: 259
Model: ARMA(0, 1) Log Likelihood 850.285
Method: css-mle S.D. of innovations 0.009
Date: Sun, 22 Jan 2017 AIC -1694.570
Time: 00:44:06 BIC -1683.899
Sample: 0 HQIC -1690.279
=====
coef std err z P>|z| [95.0% Conf. Int.]
-----
const -0.0006 0.001 -0.953 0.342 -0.002 0.001
ma.L1.y 0.1411 0.056 2.505 0.013 0.031 0.252
=====
Roots
=====
Real Imaginary Modulus Frequency
-----
MA.1 -7.0862 +0.0000j 7.0862 0.5000
=====
"""

```

```
1 In [16]: model2=arima_model.ARIMA(logReturn.values, order=(1,0,0)).fit()
```

```
RUNNING THE L-BFGS-B CODE
***
Machine precision = 2.220D-16
N =                2      M =                12
This problem is unconstrained.
At X0              0 variables are exactly at the bounds

At iterate    0      f= -3.28458D+00      |proj g|=  2.81846D-02
At iterate    5      f= -3.28458D+00      |proj g|=  2.97518D-03
At iterate   10      f= -3.28458D+00      |proj g|=  2.88435D-02
At iterate   15      f= -3.28458D+00      |proj g|=  2.36255D-05
***
Tit   = total number of iterations
Tnfx  = total number of function evaluations
Tnint = total number of segments explored during Cauchy searches
Skip  = number of BFGS updates skipped
Nact  = number of active bounds at final generalized Cauchy point
Projg = norm of the final projected gradient
F     = final function value

   N      Tit      Tnfx  Tnint  Skip  Nact      Projg      F
   2       17       20     1     0     0     0.000D+00  -3.285D+00
F = -3.2845776032338465
CONVERGENCE: NORM_OF_PROJECTED_GRADIENT_<=_PGTOL

Cauchy          time 0.000E+00 seconds.
Subspace minimization time 0.000E+00 seconds.
Line search      time 0.000E+00 seconds.

Total User time 0.000E+00 seconds.
```

```
1 In [17]: model2.summary()
```

```
2 Out[17]:
```

```
3 <class 'statsmodels.iolib.summary.Summary'>
```

```
"""
                                ARMA Model Results
=====
Dep. Variable:                  y      No. Observations:                259
Model:                        ARMA(1, 0)  Log Likelihood                850.706
Method:                      css-mle     S.D. of innovations                0.009
Date:                        Sun, 22 Jan 2017  AIC                    -1695.411
Time:                        00:44:08      BIC                    -1684.741
Sample:                      0           HQIC                    -1691.121

=====
              coef      std err          z      P>|z|      [95.0% Conf. Int.]
-----
const          -0.0006      0.001      -0.913      0.362      -0.002      0.001
ar.L1.y         0.1615      0.061       2.638      0.009       0.042      0.281

                                Roots
=====
              Real          Imaginary          Modulus          Frequency
-----
AR.1          6.1939          +0.0000j          6.1939          0.0000
=====
"""
```

我们分别构建了 MA(1) 模型和 AR(1) 模型。在这两个模型中，系数都是显著的。因此，我们无法通过系数的显著性来选择模型。这时，我们可以通过 AIC 来选择模型。很明显，*model2* 的 AIC 较小。因此，我们应该选择 *model2*，也就是 AR(1) 模型。

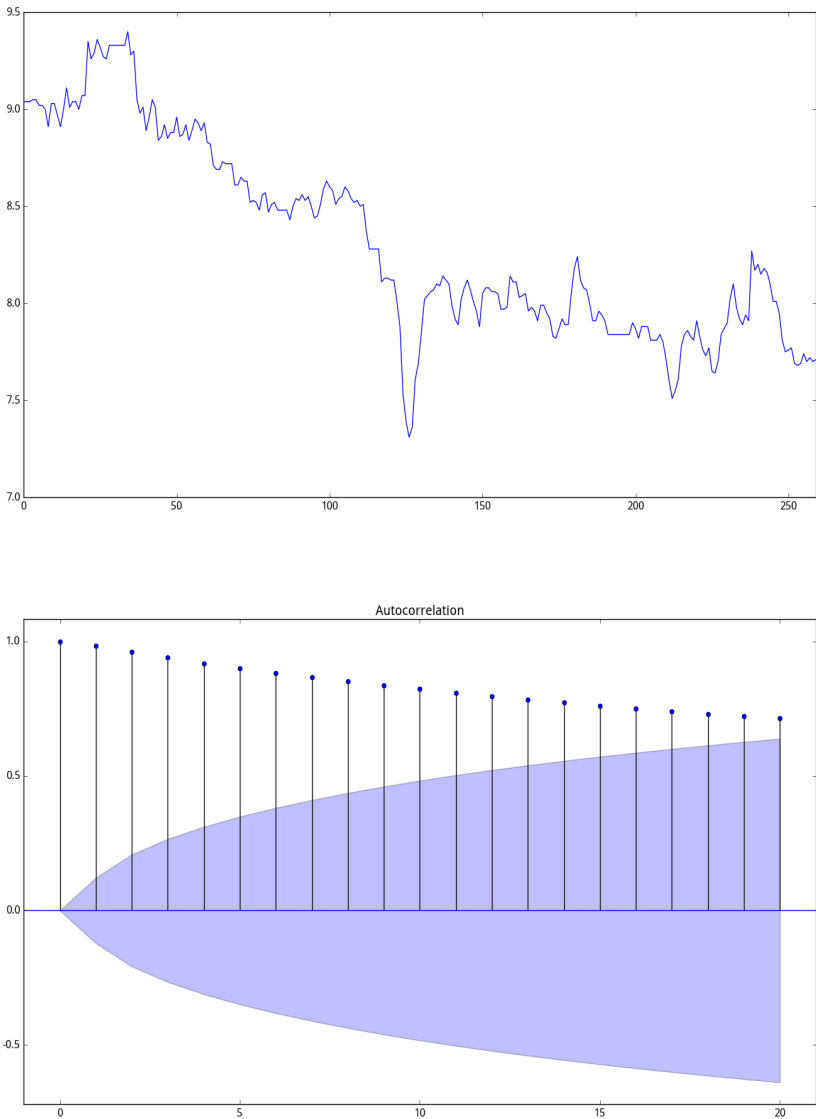


图 24.4: 题 7

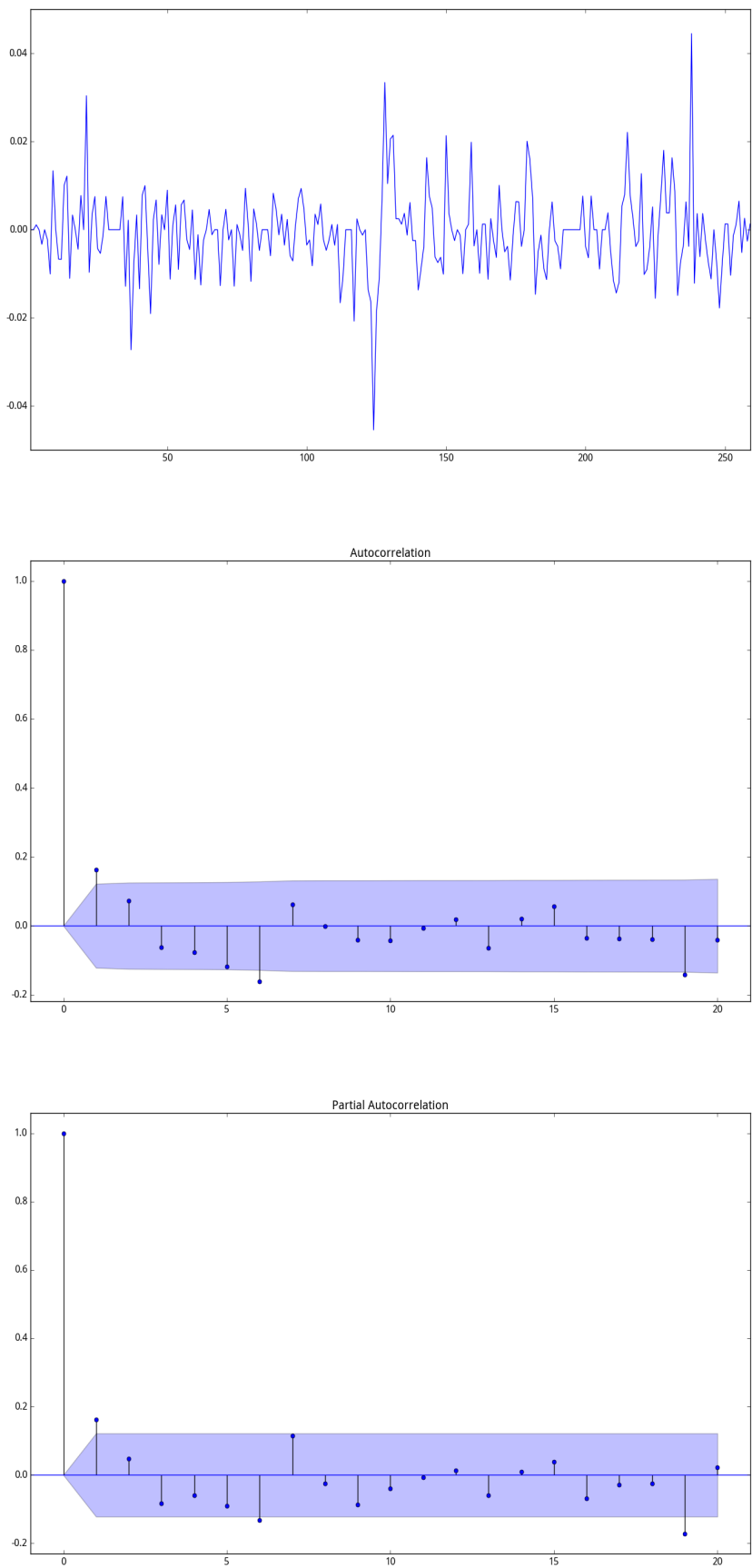


图 24.5: 题 7

```

1 In [1]: baiyun=zgsy=pd.read_csv(' baiyun.csv ',index_col= Date )
2
3 In [2]: baiyun.index=pd.to_datetime(baiyun.index)
4
5 In [3]: clprice=baiyun.Close

```

```

(a)
1 In [4]: logReturn=pd.Series((np.log(clprice))).diff().dropna()
2
3 In [5]: logReturn.plot()
4 Out[5]: <matplotlib.axes._subplots.AxesSubplot at 0x7efdb08b41d0>

```

```

1 In [6]: adf=ADF(logReturn,lags=6)
2
3 In [7]: print(adf.summary().as_text())
4     Augmented Dickey–Fuller Results
5     =====
6 Test Statistic                -6.189
7 P-value                      0.000
8 Lags                          6
9
10 Trend: Constant
11 Critical Values: -3.46 (1%), -2.87 (5%), -2.57 (10%)
12 Null Hypothesis: The process contains a unit root.
13 Alternative Hypothesis: The process is weakly stationary.

```

ADF 检验的统计量小于 5% 的临界值，因此我们能够拒绝原假设，收益率序列是平稳的。

```

(b)
1 In [8]: plot_acf(logReturn,lags=20)
2 Out[8]: <matplotlib.figure.Figure at 0x7f24e3ae7c50>
3
4 In [9]: plot_pacf(logReturn,lags=20)
5 Out[9]: <matplotlib.figure.Figure at 0x7f24e39f5ef0>
6
7 In [10]: model1=arima_model.ARIMA(logReturn.values,order=(0,0,2)).fit()
8
9 In [11]: model2=arima_model.ARIMA(logReturn.values,order=(2,0,0)).fit()
10
11 In [12]: model1.aic
12 Out[12]: -1563.422363156626
13
14 In [13]: model2.aic

```

```
15 Out[13]: -1561.4132520773128
```

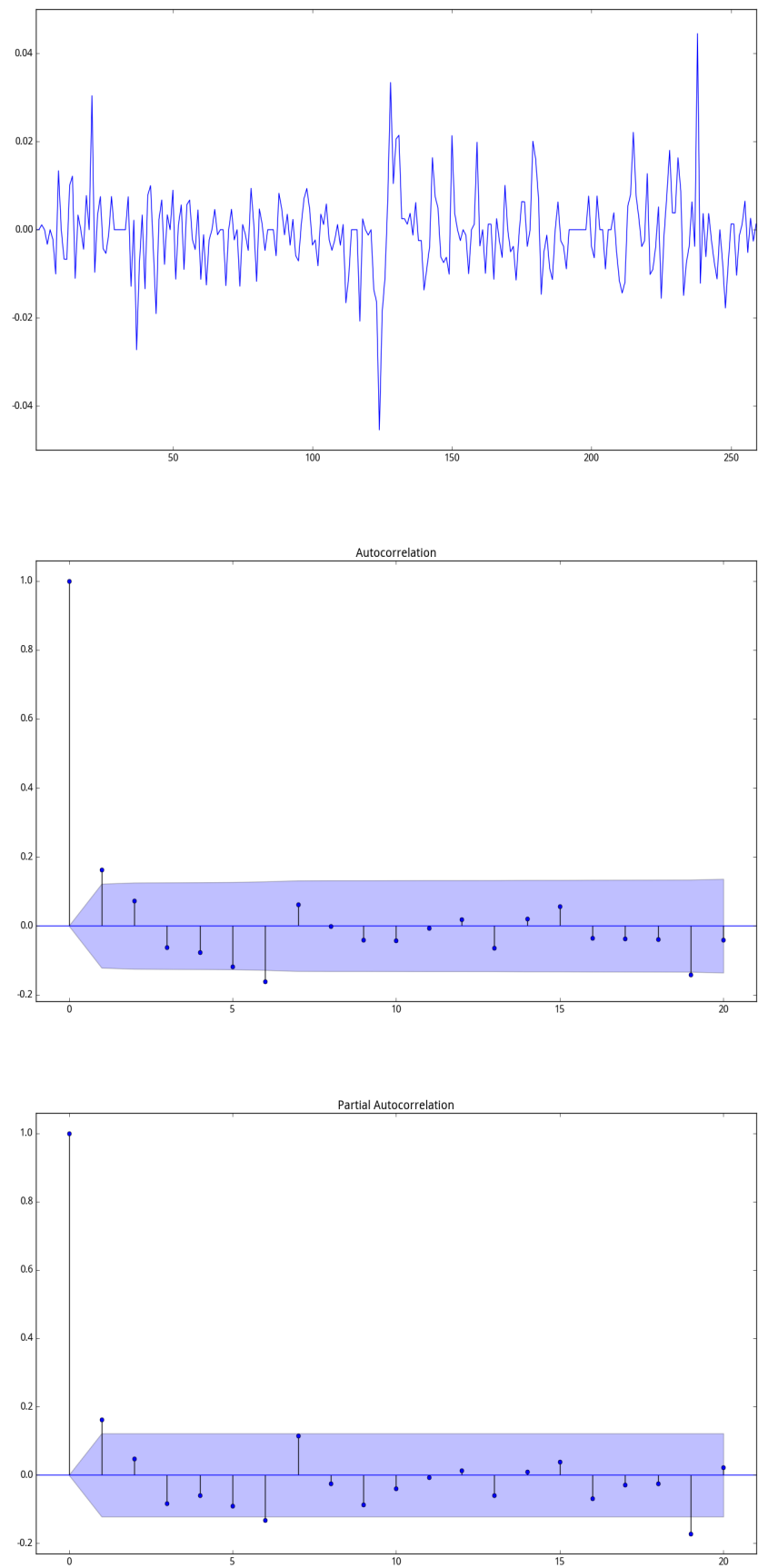


图 24.6: 题 8

在 ACF 图中，二阶的自相关系数是显著的。而在 PACF 图中，二阶的偏自相关系数也是显著的。因此，我们可以选择 MA(2) 与 AR(2) 模型。同其中，*model1* 的 AIC 较小。因此，我们应该选取 MA(2) 模型。

(d)

```
1 In [14]: import math
2
3 In [15]: stdresid=model2.resid/math.sqrt(model2.sigma2)
4
5 In [16]: stdresid.plot()
6 Out[16]: <matplotlib.axes._subplots.AxesSubplot at 0x7efdb077edd8>
7
8 In [17]: plot_acf(stdresid,lags=20)
9 Out[17]: <matplotlib.figure.Figure at 0x7f24e3b64dd8>
10
11 In [18]: LjungBox=stattools.q_stat(stattools.acf(stdresid)[1:13],len(stdresid))
12
13 In [19]: LjungBox[1][-1]
14 Out[19]: 0.21712592584238385
```

第一张图是模型的标准残差图，大部分的残差都在 ± 2 之内，这表明我们的模型对数据拟合得很好。第二张图是残差的 ACF 图，所有的自相关系数都是不显著的。因此，我们可以认为残差是一个白噪声。之后我们进行了 Ljung-Box 检验， p 值为 0.2171，表明残差序列为白噪声。因此，我们可以认为我们的模型是一个不错的模型。

(e)

```
1 In [20]: pd.Series(model2.forecast(10)[0]).plot()
2 Out[20]: <matplotlib.axes._subplots.AxesSubplot at 0x7efdb068a550>
```

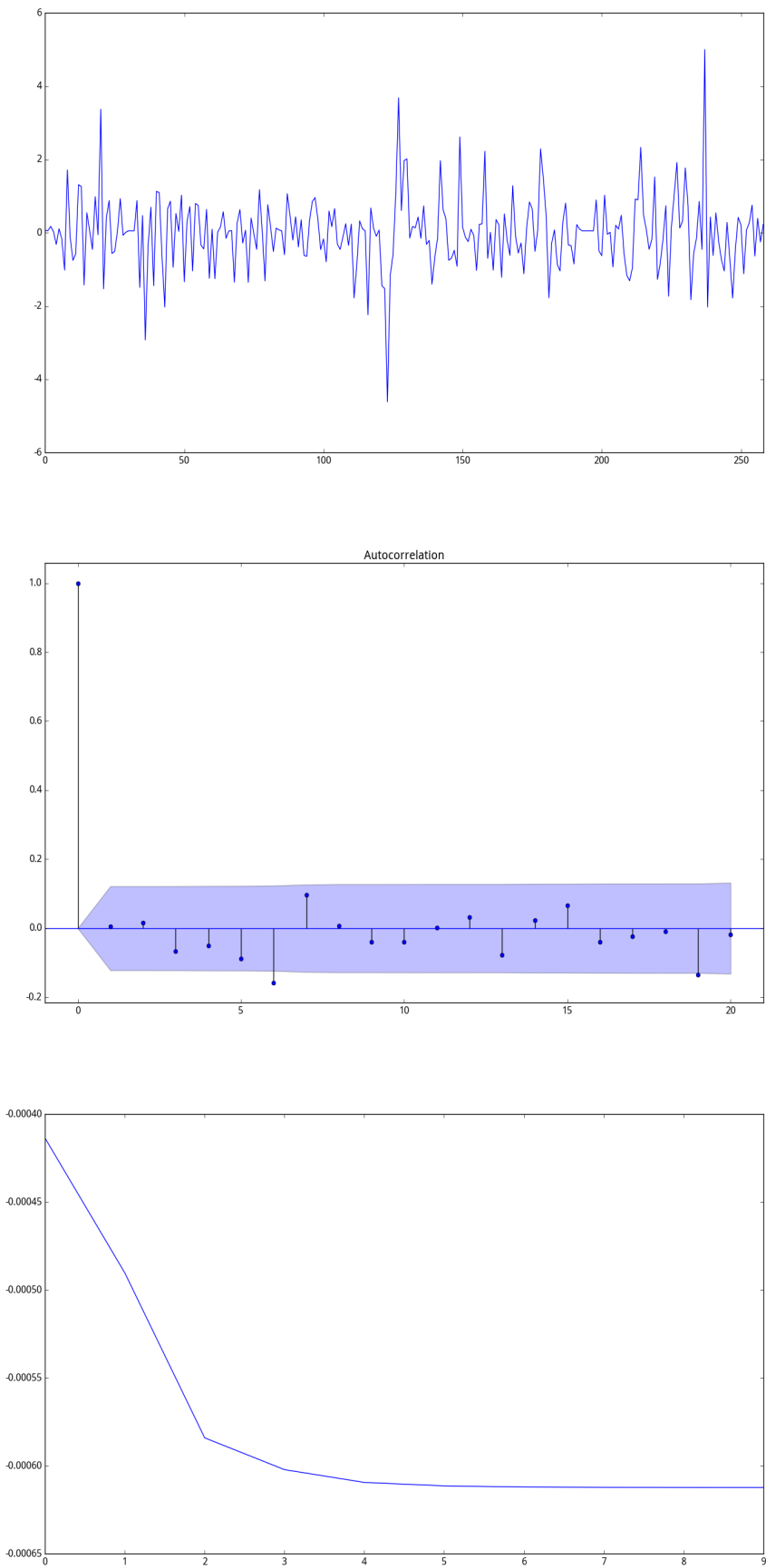


图 24.7: 题 8

第二十五章 GARCH

$$1. \because Cov(\varepsilon_t, \varepsilon_{t-k}) = E((\varepsilon_t - E(\varepsilon_t))(\varepsilon_{t-k} - E(\varepsilon_{t-k}))) = E(\varepsilon_t \varepsilon_{t-k}) - E(\varepsilon_t) E(\varepsilon_{t-k}) = E(\sigma_t^2 u_t \sigma_{t-k}^2 u_{t-k}) = E(\sigma_t^2) E(\sigma_{t-k}^2) E(u_t) E(u_{t-k}) = 0$$

$$\therefore Cor(\varepsilon_t, \varepsilon_{t-k}) = 0$$

$$2. \because \varepsilon_t = \sigma_t u_t$$

$$\therefore \varepsilon_t^2 = \sigma_t^2 u_t^2 = \sigma_t^2 + \sigma_t^2 (u_t^2 - 1) = \sigma_t^2 + \eta_t$$

$$\sigma_t^2 = \sum_{i=1}^p \alpha_i \varepsilon_t^2 \text{ 就可以转换为: } \varepsilon_t^2 = \sum_{i=1}^p \alpha_i \varepsilon_t^2 + \eta_t$$

```
3. In [1]: import pandas as pd
4. In [2]: CRSPday=pd.read_csv(' CRSPday.csv ')
5. In [3]: ibm=CRSPday.ibm
```

```
(a) In [4]: ibm.plot()
Out[4]: <matplotlib.axes._subplots.AxesSubplot at 0x7f1f4a1bd198>
```

可以看到，较大的波动后总是跟随着一些较大的波动，而较小的波动后则会跟随着较小的波动。因此，图中出现了较为明显的波动聚集现象。

```
(b) In [5]: from statsmodels.graphics.tsaplots import *
In [6]: plot_acf(ibm**2,lags=20)
```

很明显，前几阶的自相关系数都是显著的。因此，我们可以认为 *ibm* 具有 ARCH 效应。

```
(c) In [7]: from statsmodels.tsa import stattools
In [8]: LjungBox=stattools.q_stat(stattools.acf(ibm**2)[1:13],len(ibm))
In [9]: LjungBox[1][-1]
Out[9]: 8.5307836575770328e-12
```

Ljung-Box 检验的 p 值为 $8.53e-12$ ，小于 0.05。因此，该序列不是一个白噪声，我们可以认为 *ibm* 具有 ARCH 效应。

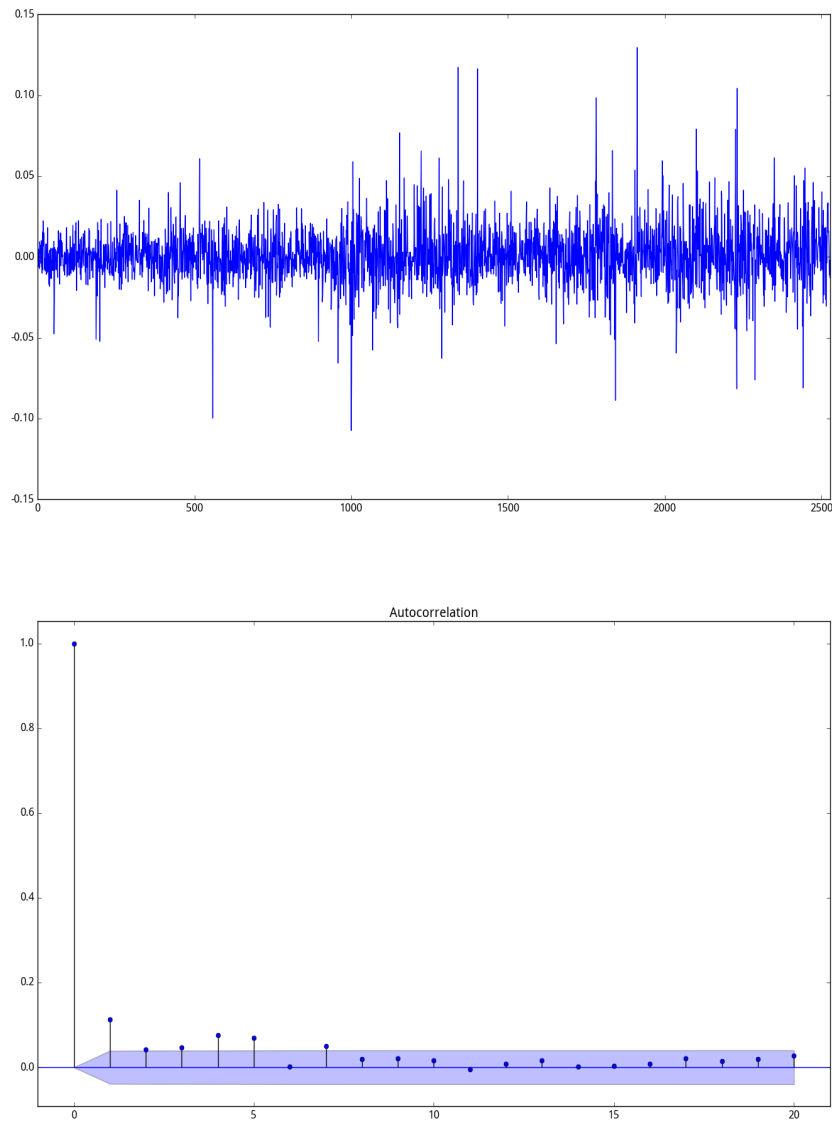


图 25.1: 题 3

```

4. In [10]: import pandas_datareader.data as web
5.         ...: import datetime as dt
6.
7. In [11]: google = web.DataReader('GOOGL', yahoo,
8.         ...:                       dt.datetime(2004,1,1),dt.datetime(2015,12,31))
9.
10. In [12]: google = google.asfreq('M', 'ffill', end)
11.
12. In [13]: googleRet = (google.Close-google.Close.shift(1))/google.Close.shift(1)
13.
14. In [14]: googleRet = googleRet.dropna()

```

(a)

```
1 In [15]: googleRet.plot()  
2 Out[15]: <matplotlib.axes._subplots.AxesSubplot at 0x7f1f487240b8>
```

从时间序列图上看, *googleRet* 并没有明显的趋势, 是平稳的时间序列。

(b)

```
1 In [16]: plot_acf(googleRet, lags=20)  
2  
3 In [17]: plot_pacf(googleRet, lags=20)
```

根据 ACF 图和 PACF 图, *googleRet* 不仅是一个平稳序列, 还是一个白噪声。

(c)

```
1 In [18]: LjungBox=stattools.q_stat(stattools.acf(googleRet)[1:13], len(  
    googleRet))  
2  
3 In [19]: LjungBox[1][-1]  
4 Out[19]: 0.78429206851907785
```

Ljung-Box 检验的 p 值为 0.7843, 大于 0.05。因此, 该序列是一个白噪声。

(d)

```
1 In [20]: (googleRet**2).plot()  
2 Out[20]: <matplotlib.axes._subplots.AxesSubplot at 0x7f1f2d3e36d8>
```

通过绘制序列平方的时间序列图, 我们可以更加容易地发现序列的波动聚集现象。在该图中, 波动聚集现象就很明显。

(e)

```
1 In [21]: plot_acf(googleRet**2, lags=20)  
2  
3 In [22]: plot_pacf(googleRet**2, lags=20)
```

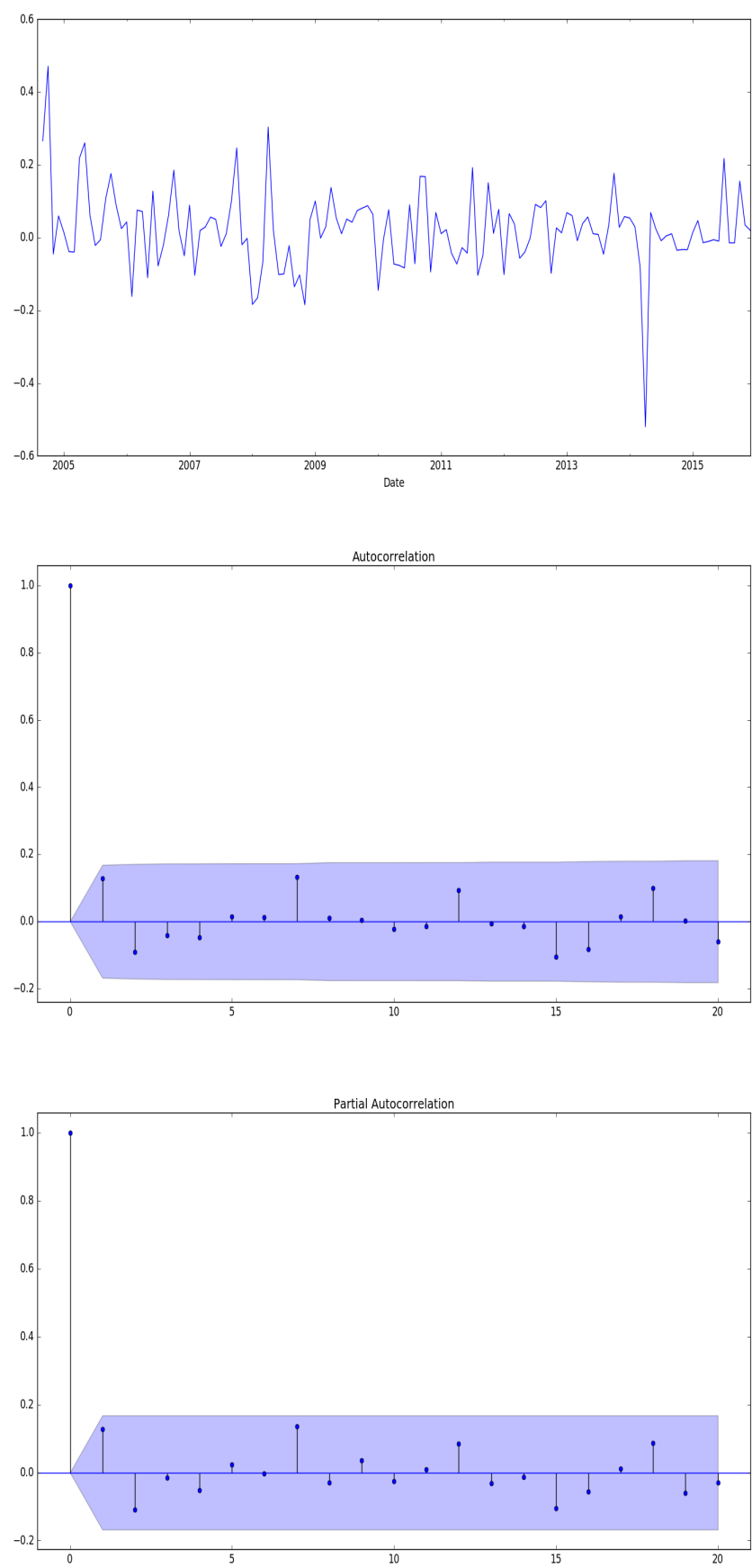


图 25.2: 题 4

(a)

```

1 In [23]: LjungBox=stattools.q_stat(stattools.acf(googleRet**2)[1:13],len(
           googleRet))
2
3 In [24]: LjungBox[1][-1]
4 Out[24]: 0.93944694696639508

```

Ljung-Box 检验的 p 值为 0.9394，大于 0.05。因此，该序列是一个白噪声，说明不存在 ARCH 效应。

(b) 我们试着拟合一个 GARCH(1,1) 模型来看看我们之前的结论是否正确。

```

1 In [25]: from arch import arch_model
2
3 In [26]: am = arch_model(googleRet)
4
5 In [27]: model = am.fit(update_freq=0)
6 Optimization terminated successfully.      (Exit mode 0)
7           Current function value: 1489.58011154
8           Iterations: 10
9           Function evaluations: 68
10          Gradient evaluations: 10
11
12 In [28]: print(model.summary())

```

```

=====
Constant Mean - GARCH Model Results
=====
Dep. Variable:          Close      R-squared:                -0.002
Mean Model:             Constant Mean  Adj. R-squared:          -0.002
Vol Model:              GARCH        Log-Likelihood:         116.703
Distribution:           Normal       AIC:                   -225.407
Method:                Maximum Likelihood  BIC:                   -213.756
                                     No. Observations:         136
Date:                  日, 1月 22 2017  Df Residuals:           132
Time:                  11:44:44      Df Model:                4
                                     Mean Model
=====
              coef    std err          t      P>|t|      95.0% Conf. Int.
-----
mu           0.0261    7.296e-03     3.580    3.430e-04    [1.182e-02, 4.042e-02]
=====
              coef    std err          t      P>|t|      95.0% Conf. Int.
-----
Volatility Model
=====
              coef    std err          t      P>|t|      95.0% Conf. Int.
-----
omega        5.6277e-03  1.698e-03     3.313    9.214e-04    [2.299e-03, 8.957e-03]
alpha[1]      0.8281      0.476        1.741    8.169e-02    [-0.104,  1.760]
beta[1]       0.0000    6.237e-02     0.000    1.000        [-0.122,  0.122]
=====
Covariance estimator: robust

```

可以看到，系数在 5% 的水平都不显著，说明 Google 的月度收益率序列确实不存在 ARCH 效应。

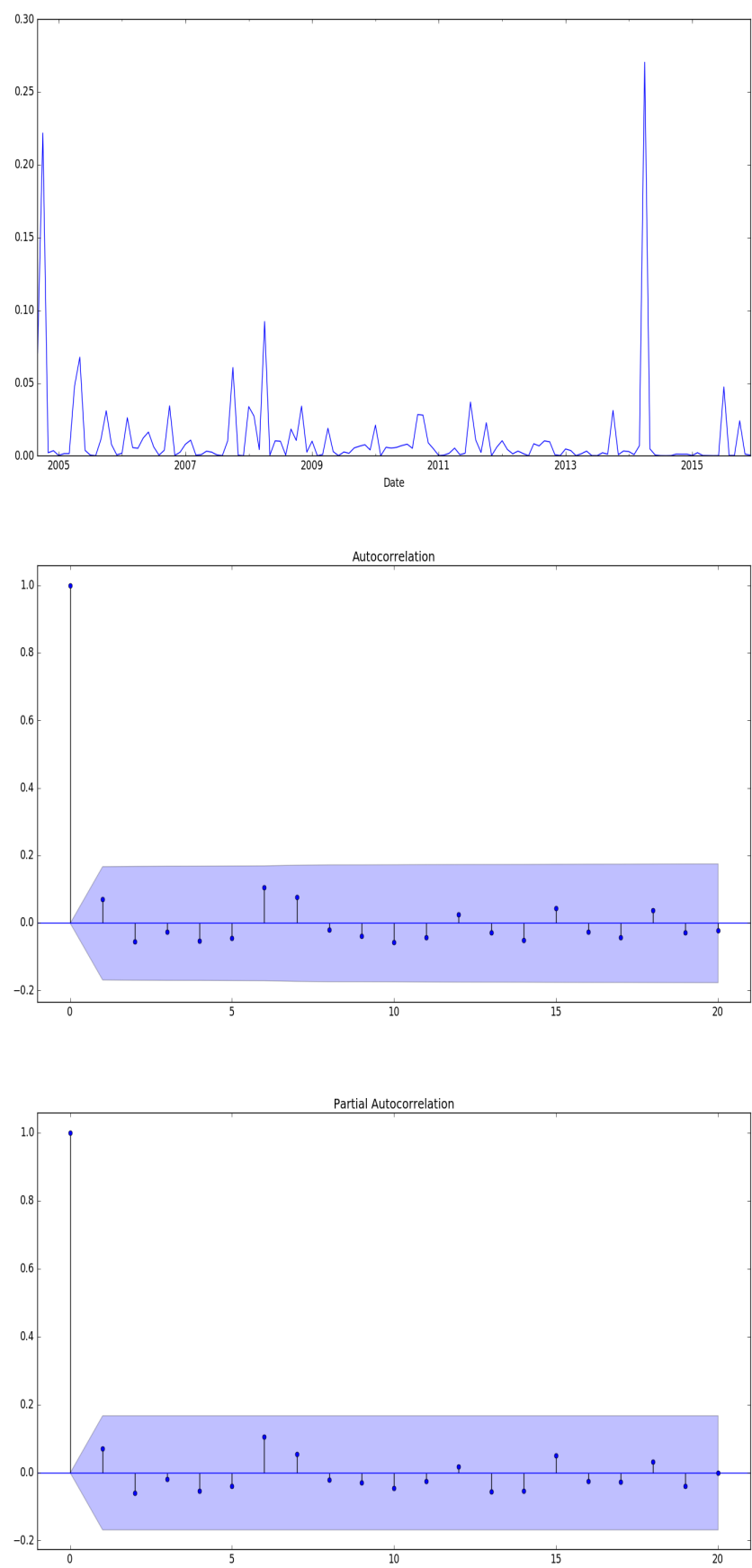


图 25.3: 题 4

```

1 In [29]: cny = web.DataReader(' CNY=X ', yahod ,
2     ...:                        dt.datetime(2015,1,1),dt.datetime(2015,12,31))
3
4 In [30]: ret = (cny.Close-cny.Close.shift(1))/cny.Close.shift(1)
5
6 In [31]: ret = ret.dropna()

```

(a)

```

1 In [25]: cny.Close.plot()
2 Out[25]: <matplotlib.axes._subplots.AxesSubplot at 0x7f7a18599ef0>

```

图中存在很明显的上升趋势，因此人民币的汇率不是一个平稳序列。

(b)

```

1 In [27]: ret.plot()
2 Out[27]: <matplotlib.axes._subplots.AxesSubplot at 0x7f7a18475630>

```

与人民币的汇率不同，其变动百分比序列并没有明显的趋势，是一个平稳序列。

(c)

```

1 In [28]: plot_acf(ret,lags=20)
2
3 In [28]: plot_pacf(ret,lags=20)

```

变动百分比序列的 ACF 图和 PACF 图中均有几个系数是显著的，因此，该序列可能并不是一个白噪声。

(d)

```

1 In [29]: LjungBox=stattools.q_stat(stattools.acf(ret)[1:13],len(ret))
2
3 In [30]: LjungBox[1][-1]
4 Out[30]: 0.0036343939189770673

```

Ljung-Box 的 p 值为 0.00363，因此我们可以拒绝原假设。该结果证明了我们上面的猜想，变动百分比序列并不是一个白噪声。

(e)

```

1 In [31]: (ret**2).plot()
2 Out[31]: <matplotlib.axes._subplots.AxesSubplot at 0x7f7a1849a550>

```

很明显，较大的波动和较小的波动分别聚集在一起。因此，变动百分比序列呈现出波动聚集现象。

(f)

```

1 In [28]: plot_acf(ret**2,lags=20)
2
3 In [28]: plot_pacf(ret**2,lags=20)

```

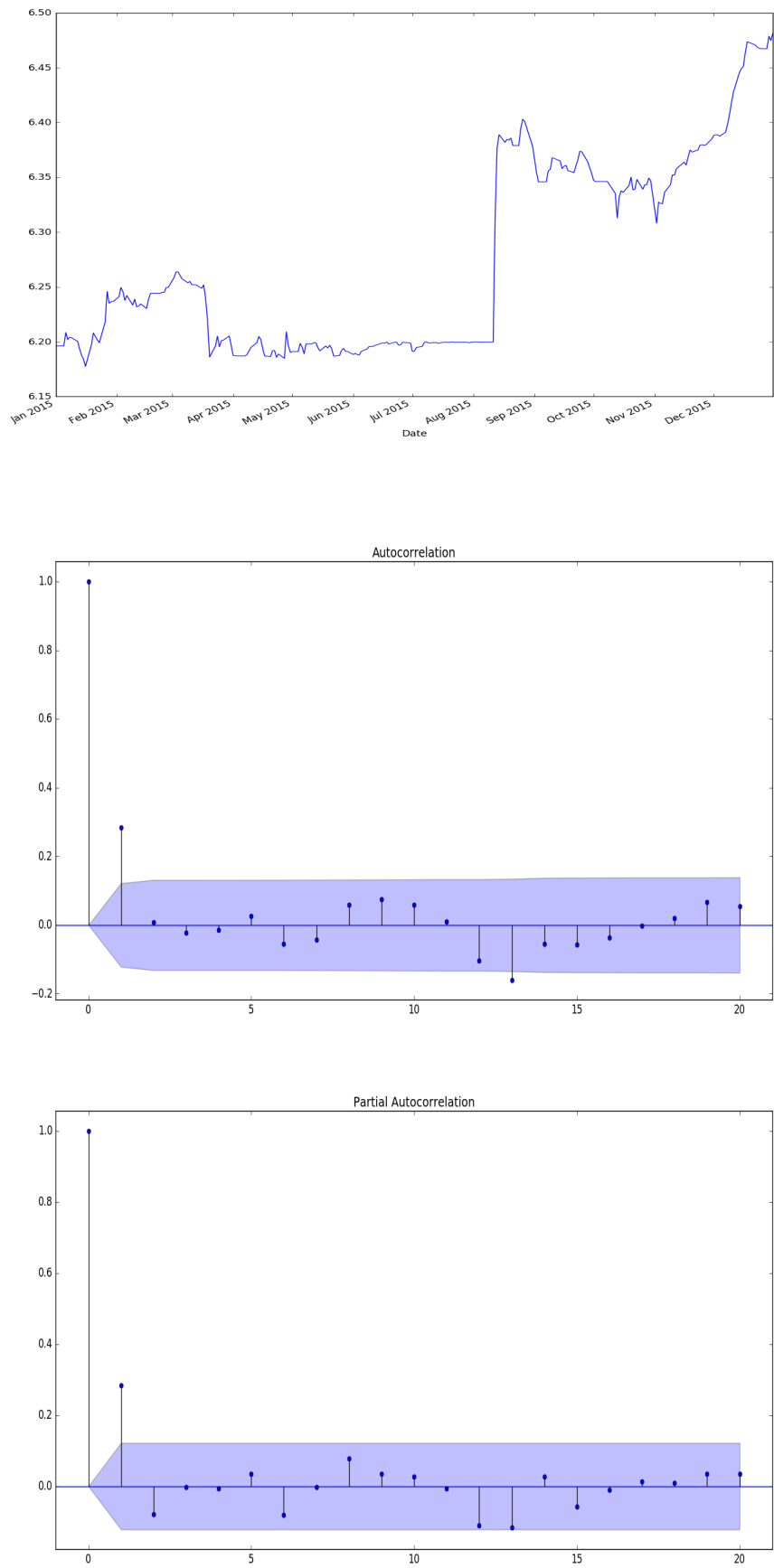


图 25.4: 题 5

```

1 In [29]: LjungBox=stattools.q_stat(stattools.acf(ret**2)[1:13],len(ret))
2
3 In [30]: LjungBox[1][-1]
4 Out[30]: 0.14677921787366507

```

Ljung-Box 的 p 值为 0.1467, 因此我们不能拒绝原假设。这与我们之前观察 ACF 图和 PACF 图得出的结论不一致, 因此, 我们可以通过拟合 GARCH 模型来看看到底有没有 ARCH 效应。

(h)

```

1 In [39]: from arch.univariate import ARX,GARCH
2
3 In [40]: model=ARX(ret,lags=1)
4
5 In [41]: model.volatility=GARCH()
6
7 In [41]: res = model.fit()
8
9 In [44]: print(res.summary())

```

```

=====
AR - GARCH Model Results
=====
Dep. Variable:          Close      R-squared:                0.081
Mean Model:              AR        Adj. R-squared:           0.077
Vol Model:               GARCH     Log-Likelihood:         1288.18
Distribution:            Normal    AIC:                   -2566.36
Method:                 Maximum Likelihood BIC:                 -2548.59
                                     No. Observations:         258
Date:                   日, 1月 22 2017   Df Residuals:           253
Time:                   19:45:35         Df Model:                5
                                     Mean Model
=====
              coef      std err          t      P>|t|      95.0% Conf. Int.
-----
Const      1.2963e-04  1.221e-05    10.619  2.444e-26  [1.057e-04,1.536e-04]
Close[1]    0.2839      0.104      2.720  6.520e-03  [7.936e-02, 0.488]
Volatility Model
=====
              coef      std err          t      P>|t|      95.0% Conf. Int.
-----
omega      1.3716e-06  1.105e-12   1.241e+06  0.000  [1.372e-06,1.372e-06]
alpha[1]    0.0500      0.218      0.229  0.819  [-0.378, 0.478]
beta[1]     0.4500      0.461      0.976  0.329  [-0.454, 1.354]
=====
Covariance estimator: robust

```

拟合的模型为

$$\begin{cases} cny_t = 1.296e - 03 + 0.284cny_{t-1} + \varepsilon_t \\ \varepsilon_t = \sigma_t u_t \\ \sigma_t^2 = 1.372e - 06 + 0.050\sigma_{t-1}^2 + 0.450\varepsilon_t^2 \end{cases}$$

根据拟合的模型, volatility model 的系数都不显著, 说明并不存在 ARCH 效应。

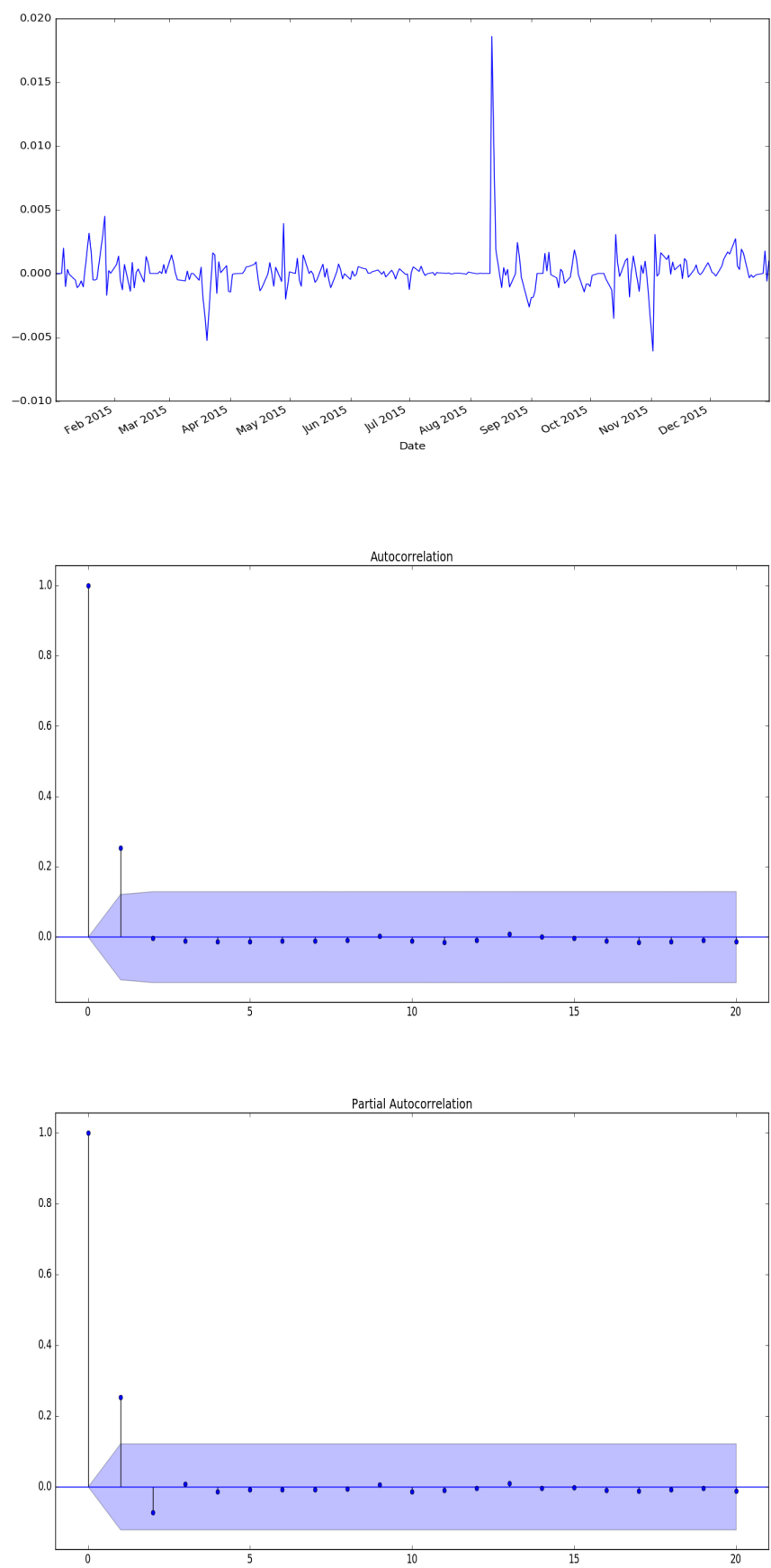


图 25.5: 题 5

第二十六章 配对交易

```
1. In [1]: import pandas as pd
2     ...:
3     ...: import pandas_datareader.data as web
4     ...:
5     ...: import datetime as dt
6     ...:
7     ...: import statsmodels.formula.api as smf
8     ...:
9     ...: import numpy as np
10    ...:
11    ...: from arch.unitroot import ADF
12
13 In [2]: zgyh = web.DataReader(' 601988.SS ', yahoo ,
14     ...:                       dt.datetime(2009,1,2),dt.datetime(2012,12,31))
15
16 In [3]: zxyh = web.DataReader(' 601998.SS ', yahoo ,
17     ...:                       dt.datetime(2009,1,2),dt.datetime(2012,12,31))
```

(a)

```
1 In [4]: train_zgyh = zgyh.Close[ 2009 : 2011 ]
2
3 In [5]: train_zxyh = zxyh.Close[ 2009 : 2011 ]
4
5 In [6]: train_zxyh.corr(train_zgyh)
6 Out[6]: 0.81995687956478036
```

两者的相关系数高达 0.82，表明我们可以对这两支股票进行配对。

(b)

```
1 In [7]: trainset=pd.concat([train_zxyh , train_zgyh ],1)
2
3 In [8]: trainset.columns=[ zxyh , zgyh ]
4
5 In [9]: trainset=np.log(trainset)
6
```

```

7 In [10]: model = smf.ols( zgyh~zxyh ,data=trainset).fit()
8
9 In [11]: resid = model.resid
10
11 In [12]: adf=ADF(resid,trend=nc)
12
13 In [13]: print(adf.summary().as_text())
14     Augmented Dickey–Fuller Results
15     =====
16 Test Statistic                -2.428
17 P-value                      0.015
18 Lags                        1
19
20 Trend: No Trend
21 Critical Values: -2.57 (1%), -1.94 (5%), -1.62 (10%)
22 Null Hypothesis: The process contains a unit root.
23 Alternative Hypothesis: The process is weakly stationary.

```

单位根检验的统计量为 -2.428 ，小于 5% 的临界值。因此，价差序列是平稳的。

(c)

```

1 In [14]: spread = trainset.zgyh - trainset.zxyh
2
3 In [15]: resid.describe()
4 Out[15]:
5 count    7.280000e+02
6 mean      2.976862e-16
7 std       6.673705e-02
8 min       -1.094785e-01
9 25%       -5.151544e-02
10 50%       -2.006203e-02
11 75%       4.786318e-02
12 max       2.227136e-01
13 Name: None, dtype: float64

```

2. 略

```

3.
1 In [16]: sh50 = pd.read_csv( 'Data/Part5/001/sh50.csv' )
2
3 In [17]: code = sh50.Code[:20]
4
5 #定义计算SSD函数
6 In [18]: def pair(a):

```

```

7      ...:      if len(a) == 2:
8      ...:          stock_a = web.DataReader(' %d.SS % a[0]', yahoo ,
9      ...:                                     dt.datetime(2015,1,1),dt.datetime(2015,6,30))
10     ...:          stock_b = web.DataReader(' %d.SS % a[1]', yahoo ,
11     ...:                                     dt.datetime(2015,1,1),dt.datetime(2015,6,30))
12     ...:          x = np.log(stock_a.Close)
13     ...:          y = np.log(stock_b.Close)
14     ...:          ret_x = x.diff()[1:]
15     ...:          ret_y = y.diff()[1:]
16     ...:          standard_x = (1+ret_x).cumprod()
17     ...:          standard_y = (1+ret_y).cumprod()
18     ...:          SSD = ((standard_x-standard_y)**2).sum()
19     ...:          return(pd.DataFrame({' SSD ':[SSD], ' stock1' :[a[0]], ' stock2' :[a
20     ...:                                [1]]}))
21     ...:      else:
22     ...:          dat = pd.DataFrame({' SSD ':[0], ' stock1' :[0], ' stock2' :[0]})
23     ...:          for i in range(1,len(a)):
24     ...:              dat = dat.append(pair([a[0],a[i]]))
25     ...:          return(dat.append(pair(a[1:])))
26     ...:
27
28 In [19]: pair_SSD = pair(code.values)
29 #选取SSD值小于1的股票进行配对
30 In [20]: pair_SSD = pair_SSD[pair_SSD.SSD!=0]
31
32 In [21]: pair_SSD = pair_SSD[pair_SSD.SSD<=1]
33 #输出结果省略

```

第五部分

技术指标

第二十七章 K 线

```
1. In [1]: %paste
2 import pandas as pd
3 import numpy as np
4 import candle as cd
5 import matplotlib.pyplot as plt
6
7 shzheng = pd.read_csv(' Data/Part5/001/problem27-1.csv ',
8                       index_col= date )
9 shzheng.index.name= Date
10 shzheng.index = pd.to_datetime(shzheng.index, format= %Y-%m-%d )
11 shzheng13=shzheng[ 2013-03-01 : 2013-05-31 ]
12 cd.candleVolume(shzheng13,
13                 candletitle= Candle Plot of Shanghai Composite Index ,
14                 bartitle= volume )
```

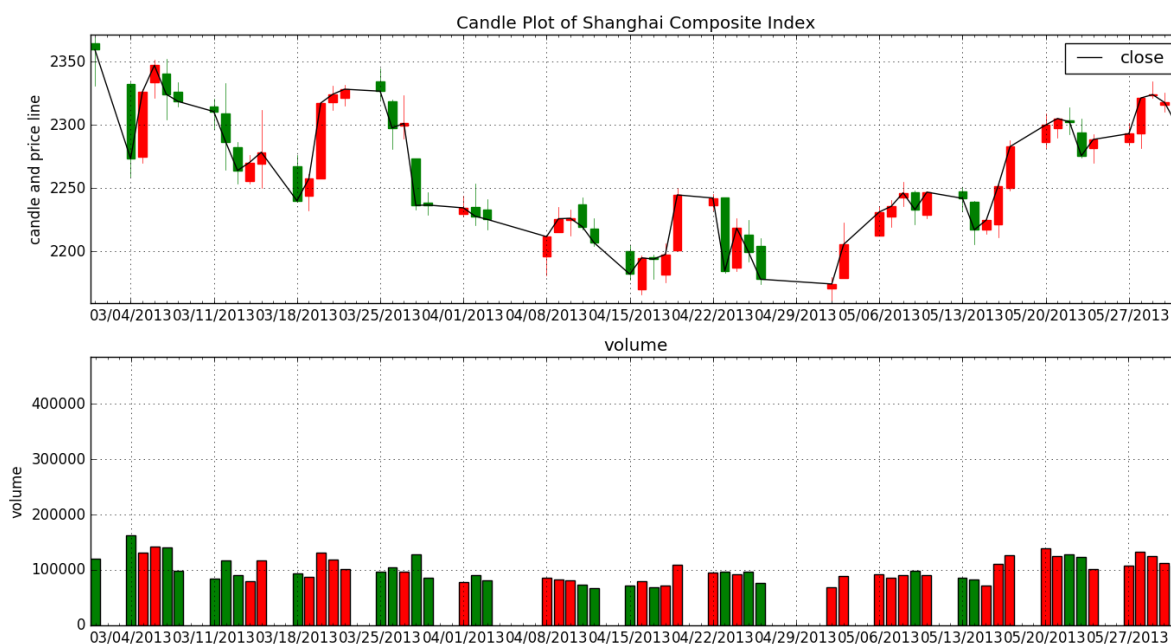


图 27.1: 题 1

```

2.
1 In [1]: shzheng = pd.read_csv(' Data/Part5/001/problem27-2.csv ',
2     ....:                      index_col= date )
3
4 In [2]: shzheng.index.name= Date
5
6 In [3]: shzheng.index = pd.to_datetime(shzheng.index,format= %Y-%m-%d )
7
8 In [4]: shzheng131=shzheng[' 2013-01-01 ' : 2013-06-30 ]
9
10 In [5]: CL_OP = shzheng131.Close - shzheng131.Open
11
12 In [6]: CL_OP.describe()
13 Out[6]:
14 count      78.000000
15 mean       -1.854376
16 std        25.289310
17 min        -105.627076
18 25%        -15.605286
19 50%         0.348023
20 75%        11.783997
21 max         59.637939
22 dtype: float64
23
24 In [7]: Doji = pd.Series(np.where(np.abs(CL_OP.values) < 8,1,0),
25     ....:                  index=CL_OP.index)
26
27 In [8]: Doji[Doji==1].index
28 Out[8]:
29 DatetimeIndex([' 2013-03-01 ', ' 2013-03-08 ', ' 2013-03-11 ', ' 2013-03-21 ',
30     ' 2013-03-22 ', ' 2013-03-27 ', ' 2013-03-29 ', ' 2013-04-01 ',
31     ' 2013-04-02 ', ' 2013-04-03 ', ' 2013-04-10 ', ' 2013-04-17 ',
32     ' 2013-04-22 ', ' 2013-05-02 ', ' 2013-05-07 ', ' 2013-05-08 ',
33     ' 2013-05-13 ', ' 2013-05-15 ', ' 2013-05-22 ', ' 2013-05-24 ',
34     ' 2013-05-27 ', ' 2013-05-29 ', ' 2013-05-30 ', ' 2013-06-03 ',
35     ' 2013-06-05 ', ' 2013-06-18 ', ' 2013-06-19 ', ' 2013-06-26 ',
36     ' 2013-06-27 ' ],
37     dtype= datetime64[ns] , name= Date , freq=None)
38
39 In [9]: cd.candlePlot(shzheng131,
40     ....:               title= Candle Plot of Shanghai Composite Index )

```

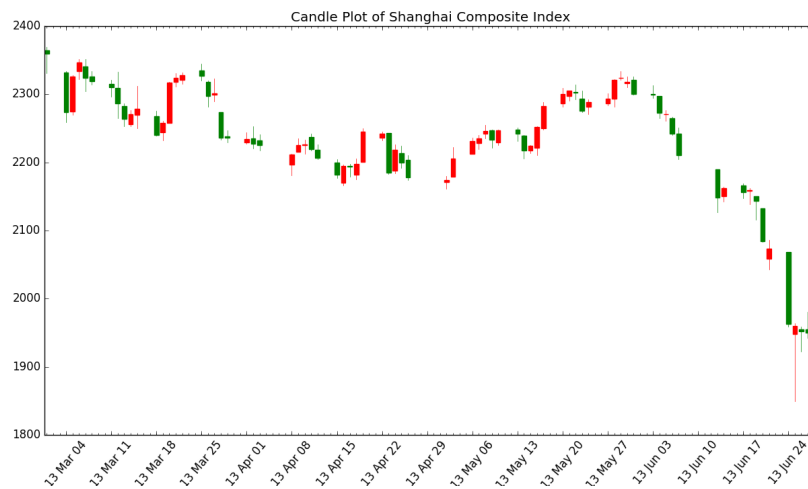


图 27.2: 题 2

3.

(a)

```

1 In [1]: shzheng = pd.read_csv( Data/Part5/001/problem27-3.csv ,
2     ....:                      index_col= date )
3
4 In [2]: shzheng.index.name= Date
5
6 In [3]: shzheng.index = pd.to_datetime(shzheng.index, format= %Y-%m-%d )

```

```

1 #刻画蜡烛实体
2 In [4]: CL_OP = shzheng.Close - shzheng.Open
3
4 In [5]: CL_OP.describe()
5 Out[5]:
6 count      243.000000
7 mean         2.244368
8 std         22.884043
9 min        -73.685059
10 25%        -11.442016
11 50%         1.513183
12 75%         14.296997
13 max         91.843994
14 dtype: float64
15
16 In [6]: dat = pd.concat([CL_OP,CL_OP.shift(1),CL_OP.shift(2)],1)
17
18 In [22]: candle = np.all([np.abs(dat.iloc[:,1])<=3,dat.iloc[:,2]>11,
19     ....:                  dat.iloc[:,0]<-6,

```

```

20     ....:             np.abs(dat.iloc[:,0])<np.abs(dat.iloc[:,2])/2),
21     ....:             0)
22
23 #刻画十字星
24 In [8]: dataCOP = pd.concat([shzheng.Close,shzheng.Open.shift(1),
25     ....: shzheng.Close.shift(1),shzheng.Open.shift(2)],1)
26
27
28 In [9]: Doji = np.all([dataCOP.iloc[:,1]>dataCOP.iloc[:,3],
29     ....: dataCOP.iloc[:,1]>dataCOP.iloc[:,0],
30     ....: dataCOP.iloc[:,2]>dataCOP.iloc[:,3],
31     ....: dataCOP.iloc[:,2]>dataCOP.iloc[:,0]],0)
32
33 #刻画上涨趋势
34 In [10]: ret = (shzheng.Close-shzheng.Close.shift(1))/shzheng.Close.shift(1)
35
36 In [11]: trend = np.all([ret.shift(2)>0,ret.shift(1)>0],0)

```

(b)

```

1 In [12]: signal = np.all([candle,Doji,trend],0)
2
3 In [13]: sum(signal)
4 Out[13]: 0

```

最终，没有发现“黄昏之星”。

4.

(a)

```

1 In [28]: shzheng = pd.read_csv('Data/Part5/001/problem27-4.csv',
2     ....:             index_col= 'date' )
3
4 In [29]: shzheng.index.name= 'Date'
5
6 In [30]: shzheng.index = pd.to_datetime(shzheng.index,format= '%Y-%m-%d' )

```

```

1 #刻画蜡烛实体
2 In [31]: CL_OP = shzheng.Close - shzheng.Open
3
4 In [32]: candle = np.all([CL_OP<0,CL_OP.shift(1)>0],0)
5
6 #刻画价格差距
7 In [33]: diffPrice = np.all([shzheng.Open>shzheng.Close.shift(1),shzheng.
    Close<shzheng.Open.shift(1)],0)

```

```

8
9 #刻画上涨趋势
10 In [34]: ret = (shzheng.Close-shzheng.Close.shift(1))/shzheng.Close.shift(1)
11
12 In [35]: trend = np.all([ret.shift(2)>0,ret.shift(1)>0],0)

```

(a)

```

1 In [36]: signal = np.all([candle,diffPrice,trend],0)
2
3 In [37]: sum(signal)
4 Out[37]: 0

```

最终，我们并没有发现“看跌吞没”。

5.

(a)

```

1 #接上题
2 #刻画蜡烛实体
3 In [38]: CL_OP.describe()
4 Out[38]:
5 count      245.000000
6 mean        0.886421
7 std         28.573461
8 min         -85.802002
9 25%         -15.752930
10 50%          0.811035
11 75%         17.271973
12 max          84.511963
13 dtype: float64
14
15 #刻画蜡烛实体
16 In [39]: candle = np.all([CL_OP<0,CL_OP.shift(1)>17],0)
17
18 #刻画价格差距
19 In [40]: diffPrice = np.all([shzheng.Close<shzheng.Open.shift(1)],0)
20
21 #刻画上涨趋势
22 In [41]: trend = np.all([ret.shift(1)>0,ret.shift(2)>0],0)

```

```

1 In [42]: signal = pd.Series(np.all([candle,diffPrice,trend],0),index=CL_OP.
      index)
2

```

```
3 In [43]: signal[signal==True]
4 Out[43]:
5 Date
6 2011-01-20    True
7 2011-05-04    True
8 2011-09-22    True
9 dtype: bool
10
11 In [44]: cd.candlePlot(shzheng[ 2011-01-01 : 2011-02-15 ], 2011-01-20 )
12
13 In [45]: cd.candlePlot(shzheng[ 2011-04-15 : 2011-05-15 ], 2011-05-04 )
14
15 In [46]: cd.candlePlot(shzheng[ 2011-09-15 : 2011-10-15 ], 2011-09-22 )
```

从 K 线图的走势我们可以看出，1 月 20 号的“倾盆大雨”是虚假信号，而另外两个都是真实的信号。

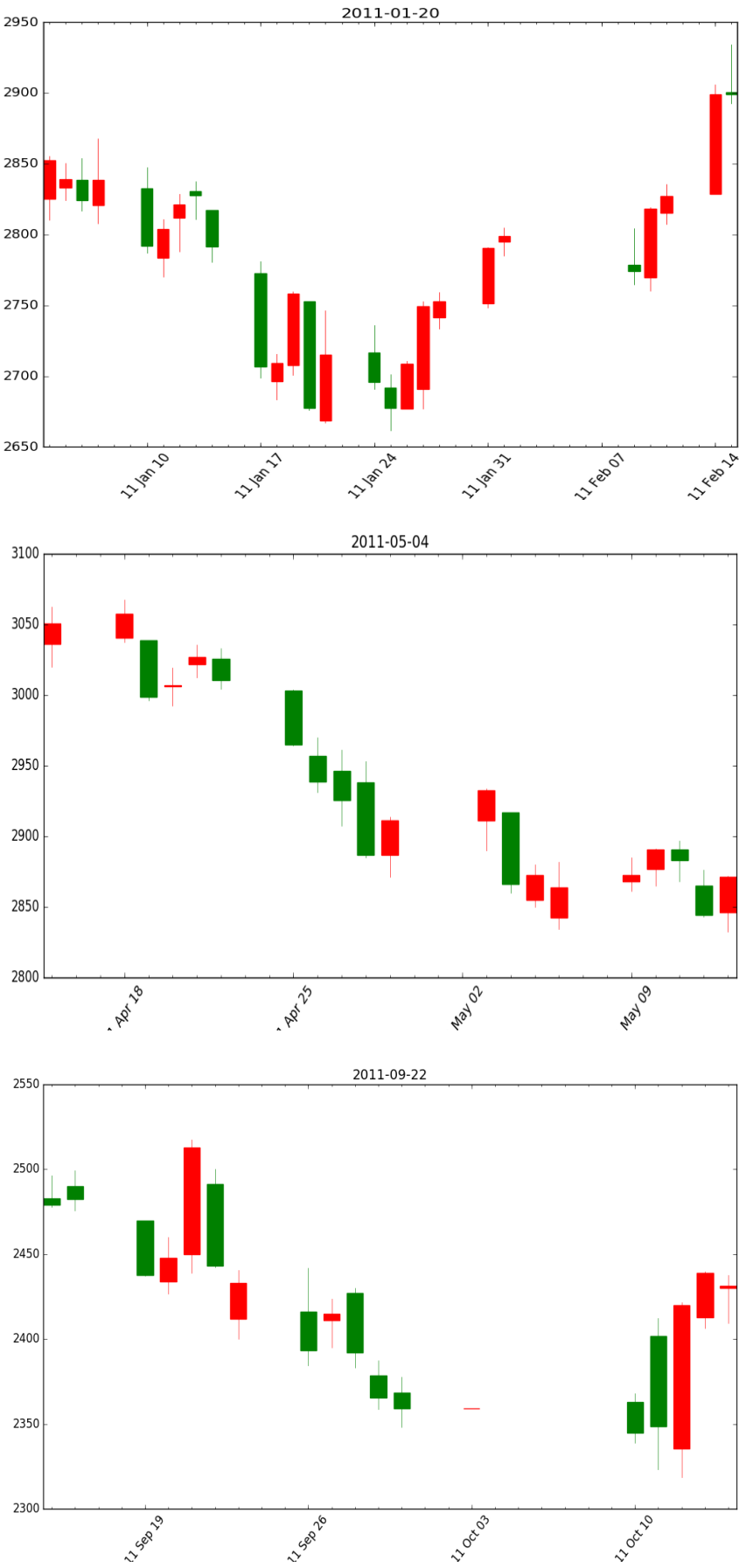


图 27.3: 题 5

第二十八章 动量

(b)

```
1 In [1]: import pandas as pd
2     ...:
3     ...: import numpy as np
4     ...:
5     ...:
6 In [2]: pufa = pd.read_csv('Data/Part5/002/problem28-1.csv', index_col='date')
7
8 In [3]: pufa.index.name = 'Date'
9
10 In [4]: pufa.index = pd.to_datetime(pufa.index, format='%Y-%m-%d')
```

(a)

```
1 #6日 动量
2 In [5]: mom6_1 = pufa.Close - pufa.Close.shift(6)
3
4 In [6]: mom6_2 = (pufa.Close - pufa.Close.shift(6))/pufa.Close.shift(6)
5
6 #30日 动量
7 In [7]: mom30_1 = pufa.Close - pufa.Close.shift(30)
8
9 In [8]: mom30_2 = (pufa.Close - pufa.Close.shift(30))/pufa.Close.shift(30)
10
11 #90日 动量
12 In [9]: mom90_1 = pufa.Close - pufa.Close.shift(90)
13
14 In [10]: mom90_2 = (pufa.Close - pufa.Close.shift(90))/pufa.Close.shift(90)
```

```
1 #定义计算交易信号函数
2 In [11]: def getSignal(x):
3     ...:     signal = np.where(x>0,1,np.where(x<0,-1,0))
4     ...:     return(signal)
5     ...:
6     ...:
```

```
7
8 #计算交易信号
9 In [12]: signal6 = getSignal(mom6_2)
10
11 In [13]: signal30 = getSignal(mom30_2)
12
13 In [14]: signal90 = getSignal(mom90_2)
14
15 #计算收益率
16 In [15]: ret = (pufa.Close-pufa.Close.shift(1))/pufa.Close.shift(1)
17
18 In [16]: ret6 = ret[6:]
19
20 In [17]: ret30 = ret[30:]
21
22 In [18]: ret90 = ret[90:]
23
24 In [19]: signal6 = pd.Series(signal6[6:], index=ret6.index)
25
26 In [20]: Mom6Ret = ret6[1:] * signal6.shift(1)[1:]
27
28 In [21]: signal30 = pd.Series(signal6[30:], index=ret30.index)
29
30 In [22]: Mom30Ret = ret30[1:] * signal30.shift(1)[1:]
31
32 In [23]: signal90 = pd.Series(signal6[90:], index=ret90.index)
33
34 In [24]: Mom90Ret = ret90[1:] * signal90.shift(1)[1:]
35 #计算获胜率
36 In [25]: winrate6 = sum(Mom6Ret>0)/sum(Mom6Ret!=0)
37
38 In [26]: winrate6
39 Out[26]: 0.44642857142857145
40
41 In [27]: winrate30 = sum(Mom30Ret>0)/sum(Mom30Ret!=0)
42
43 In [28]: winrate30
44 Out[28]: 0.44390243902439025
45
46 In [29]: winrate90 = sum(Mom90Ret>0)/sum(Mom90Ret!=0)
```

```

47
48 In [30]: winrate90
49 Out[30]: 0.44966442953020136

```

在这三个策略中，三种动量的获胜率几乎相同，使用 90 日动量的策略稍微略胜一筹。

```

1 #定义计算交易信号函数
2 In [31]: def getSignal(x):
3     ...:     first = x.iloc[:,0]
4     ...:     second = x.iloc[:,1]
5     ...:     signal = [1 if first[i]>0 and second[i]>0 else -1 if first[i]<0 and
6     ...:     second[i]<0 else 0 for i in range(len(first))]
7     ...:     return(signal)
8     ...:
9
10 #计算交易信号
11 In [32]: momen = pd.concat([mom6_2,mom90_2],1)
12
13 In [33]: momen = momen.dropna()
14
15 In [34]: signal = getSignal(momen)
16
17 In [35]: signal = pd.Series(signal,index= mom90_2.dropna().index)
18 #计算获胜率
19 In [36]: MomRet = ret90[2:] * signal.shift(2)[2:]
20
21 In [37]: sum(MomRet>0)/sum(MomRet!=0)
22 Out[37]: 0.45000000000000001

```

(a)

```

(a)
1 In [1]: import pandas as pd
2
3 In [2]: import numpy as np
4
5 In [3]: prices = pd.read_csv('Data/Part5/002/problem28-3.csv',
6     ...:                     index_col='date')
7
8 In [4]: prices.index.name='Date'
9

```

```

10 In [5]: prices.index = pd.to_datetime(prices.index,format= '%Y-%m-%d' )
11
12 In [6]: prices = prices.asfreq( M ,how= end ,method= ffill' )
13
14 In [7]: ret_pf = (prices.pfyh - prices.pfyh.shift(3))/prices.pfyh.shift(3)
15
16 In [8]: ret_zg = (prices.zgyh - prices.zgyh.shift(3))/prices.zgyh.shift(3)
17
18 In [9]: ret_ms = (prices.msyh - prices.msyh.shift(3))/prices.msyh.shift(3)
19
20 In [10]: ret_gs = (prices.gsyh - prices.gsyh.shift(3))/prices.gsyh.shift(3)
21
22 In [11]: ret_js = (prices.jsyh - prices.jsyh.shift(3))/prices.jsyh.shift(3)

```

```

1 In [12]: %paste
2 cash = np.ones(len(ret_pf))*1000000
3 stock = np.zeros(len(ret_pf))
4
5 for i in range(4,len(cash)):
6     momen = pd.Series([ret_gs[i-1],ret_js[i-1],
7         ret_ms[i-1],ret_pf[i-1],ret_zg[i-1]])
8     if i!=4:
9         momen_lag = pd.Series([ret_gs[i-2],ret_js[i-2],
10         ret_ms[i-2],ret_pf[i-2],ret_zg[i-2]])
11         cash[i] = cash[i-1] + 10000 * prices.iloc[:,momen_lag.idxmax()][i]
12         stock[i] = 10000 * prices.iloc[:,momen.idxmax()][i]
13         cash[i] -= 10000 * prices.iloc[:,momen.idxmax()][i]
14
15 (cash[35]+stock[35]-1000000)/1000000*len(cash)
16 ## — End pasted text —
17 Out[16]: 0.072000000000000008

```

第二十九章 RSI

```
1 In [1]: %paste
2 import pandas as pd
3 import numpy as np
4 def rsi(price, period=6):
5     import pandas as pd
6     clprcChange=price-price.shift(1)
7     clprcChange=clprcChange.dropna()
8     indexprc=clprcChange.index
9     upPrc=pd.Series(0, index=indexprc)
10    upPrc[clprcChange>0]=clprcChange[clprcChange>0]
11    downPrc=pd.Series(0, index=indexprc)
12    downPrc[clprcChange<0]=-clprcChange[clprcChange<0]
13    rsidata=pd.concat([price, clprcChange, upPrc, downPrc], \
14    axis=1)
15    rsidata.columns=[ 'price' , 'PrcChange' , 'upPrc' , 'downPrc' ]
16    rsidata=rsidata.dropna();
17    SMUP=[]
18    SMDOWN=[]
19    for i in range(period, len(upPrc)+1):
20        SMUP.append(np.mean(upPrc.values[(i-period):i], \
21        dtype=np.float32))
22        SMDOWN.append(np.mean(downPrc.values[(i-period):i], \
23        dtype=np.float32))
24        rsi=[100*SMUP[i]/(SMUP[i]+SMDOWN[i]) \
25        for i in range(0, len(SMUP))]
26    indexRsi=indexprc[(period-1):]
27    rsi=pd.Series(rsi, index=indexRsi)
28    return(rsi)
29
30 ## — End pasted text —
```

(b)

```
1 In [1]: import pandas as pd
```

```

2     ...:
3     ...: import numpy as np
4
5 In [2]: jtyh = pd.read_csv(' Data/Part5/003/problem29-1.csv' ,index_col= date )
6
7 In [3]: jtyh.index.name= Date
8
9 In [4]: jtyh.index = pd.to_datetime(jtyh.index, format= %Y-%m-%d )
10
11 In [5]: RSI5 = rsi(jtyh.Close,5)
12
13 In [6]: RSI6 = rsi(jtyh.Close,6)
14
15 In [7]: RSI7 = rsi(jtyh.Close,7)
16
17 In [8]: RSI8 = rsi(jtyh.Close,8)
18
19 In [9]: RSI9 = rsi(jtyh.Close,9)
20
21 In [10]: RSI10 = rsi(jtyh.Close,10)
22
23 In [11]: RSI12 = rsi(jtyh.Close,12)
24
25 In [12]: RSI_all =pd.concat([RSI5,RSI6,RSI7,RSI8,RSI9,RSI10,RSI12],1)

```

```

1 In [13]: jtyh = pd.read_csv(' Data/Part5/003/problem29-2.csv' ,index_col= date )
2
3 In [14]: jtyh.index.name= Date
4
5 In [15]: jtyh.index = pd.to_datetime(jtyh.index, format= %Y-%m-%d )

```

2. (a)

```

1 In [16]: import candle
2
3 In [17]: candle.candlePlot(jtyh,' Candle Plot of Bank of Communications ')

```

```

1 In [18]: RSI6 = rsi(jtyh.Close,6)
2
3 In [19]: RSI30 = rsi(jtyh.Close,30)
4
5 In [20]: RSI = pd.concat([RSI6,RSI30],1)

```

```

6
7 In [21]: RSI.columns=[ RSI6 ; RSI30 ]
8
9 In [22]: RSI.plot()

```

(a)

```

1 In [23]: signal = np.where(RSI6>90,-1,np.where(RSI6<10,1,0))
2
3 In [24]: ret = (jtyh.Close-jtyh.Close.shift(1))/jtyh.Close.shift(1)
4
5 In [25]: ret6 = ret[6:]
6
7 In [26]: signal = pd.Series(signal,index=ret6.index)
8
9 In [27]: RSI_ret = ret6[1:] * signal.shift(1)[1:]
10
11 In [28]: sum(RSI_ret>0)/sum(signal!=0)
12 Out[28]: 0.40566037735849059

```

```

1 In [29]: RSI_all = pd.concat([RSI6,RSI30,RSI6.shift(1),RSI30.shift(1)],1)
2
3 In [30]: RSI_all = RSI_all.dropna()
4
5 In [31]: RSI_all.columns = [ RSI6 ; RSI30 ; L_RSI6 ; L_RSI30 ]
6
7 In [32]: signal = [1 if RSI_all.RSI6[i] > RSI_all.RSI30[i] and RSI_all.L_RSI6
8     ...:             -1 if RSI_all.RSI6[i] < RSI_all.RSI30[i] and RSI_all.
9     ...:             L_RSI6[i] > RSI_all.L_RSI30[i] else 0 \
10     ...:             for i in range(len(RSI_all))]
11
12 In [33]: ret30 = ret[30:]
13
14 In [34]: signal = pd.Series(signal,index=ret30.index)
15
16 In [35]: RSI_ret = ret30[1:] * signal.shift(1)[1:]
17
18 In [36]: sum(RSI_ret>0)/sum(signal!=0)
19 Out[36]: 0.33695652173913043

```

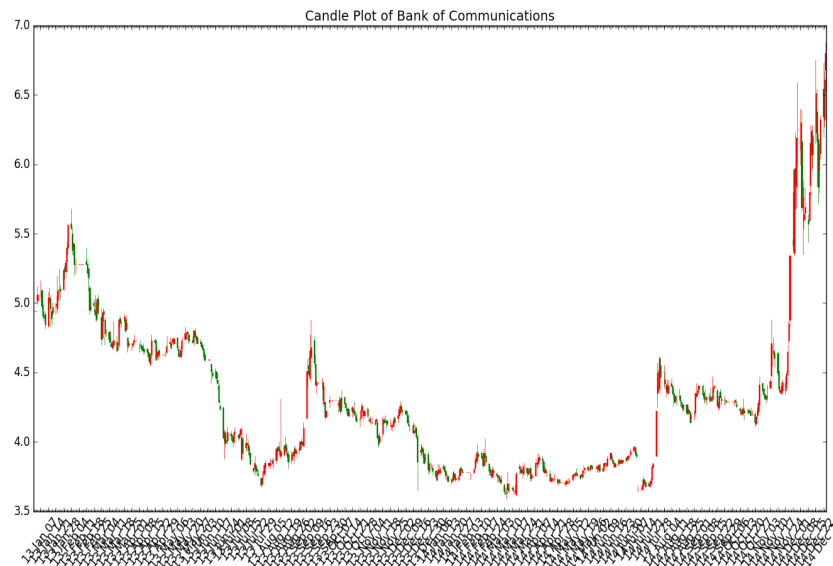


图 29.1: 题 2

(3)

```

1 In [37]: RSI_all = pd.concat([RSI6,RSI30],1)[24:]
2
3 In [38]: signal = np.zeros(len(RSI_all))
4
5 In [39]: RSI_all.columns = [ 'RSI6' , 'RSI30' ]
6 #计算交易信号
7 In [40]: for i in range(3,len(RSI_all)):
8     ...:     if RSI_all.RSI6[i-3] < RSI_all.RSI30[i-3] and \
9     ...:     RSI_all.RSI6[i-2] > RSI_all.RSI30[i-2] and \
10    ...:     RSI_all.RSI6[i-1] > RSI_all.RSI30[i-1] and \
11    ...:     (RSI_all.RSI6[i-1] - RSI_all.RSI30[i-1]) < (RSI_all.RSI6[i-2] -
    ...:     RSI_all.RSI30[i-2]) and \
12    ...:     (RSI_all.RSI6[i] - RSI_all.RSI30[i]) > (RSI_all.RSI6[i-1] - RSI_all.
    ...:     RSI30[i-1]):
13    ...:         signal[i] = 1
14    ...:
15    ...:
16
17 In [41]: signal = pd.Series(signal,index = RSI_all.index)
18
19 In [42]: signal[signal==1]
20 Out[42]:
21 Date
22 2013-08-01    1.0
23 2013-10-29    1.0

```

```
24 2014-01-09    1.0
25 2014-04-10    1.0
26 2014-09-04    1.0
27 2014-11-13    1.0
28 dtype: float64
```

鉴于买点较少，我们可以直接查到买点次日的收益率大于 0。

```
1 In [33]: buyDate=signal[signal==1]
2
3 In [34]: ret.shift(-1)[buyDate.index]
4 Out[34]:
5 Date
6 2013-08-01    -0.002611
7 2013-10-29     0.007092
8 2014-01-09     0.002639
9 2014-04-10     0.002577
10 2014-09-04     0.006865
11 2014-11-13    -0.008621
12 Name: Close, dtype: float64
```

第三十章 均线

```
1. In [1]: import pandas as pd
2         ...: import numpy as np
3         ...: import movingAverage as ma
4         ...:
5
6 In [2]: dow = pd.read_csv('Data/Part5/004/problem30-1.csv', index_col='date')
7
8 In [3]: dow.index.name = 'Date'
9
10 In [4]: dow.index = pd.to_datetime(dow.index, format='%Y-%m-%d')
11
12 In [5]: sma = ma.smaCal(dow.Close, 30)
13
14 In [6]: wma = ma.wmaCal(dow.Close, [(i+1)/465 for i in range(30)])
15
16 In [7]: ema = ma.ewmaCal(dow.Close, 30, 0.8)
17
18 In [8]: maValues = pd.concat([sma, wma, ema], 1).replace(0, np.nan).dropna()
19
20 In [9]: maValues.columns = ['sma', 'wma', 'ema']
21
22 In [10]: maValues.plot()
```

可以看到，三者是非常接近的。读者可以自行更改 WMA 和 EMA 的权重，观察发生的变化。

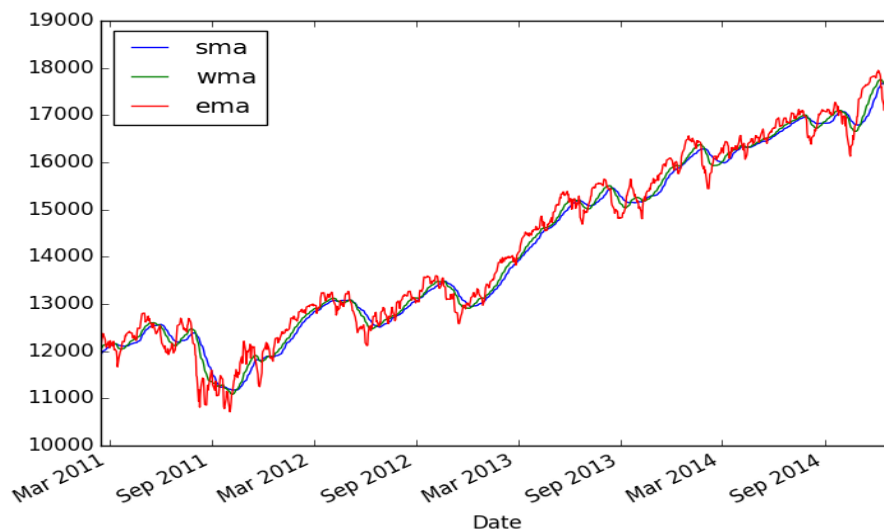


图 30.1: 题 1

```

2. In [11]: zgyh = pd.read_csv('Data/Part5/004/problem30-2.csv', index_col='date')
3. In [12]: zgyh.index.name = 'Date'
4.
5. In [13]: zgyh.index = pd.to_datetime(zgyh.index, format='%Y-%m-%d')

```

```

(a) In [14]: import matplotlib.pyplot as plt
6.
7. In [15]: mom14 = (zgyh.Close - zgyh.Close.shift(14)) / zgyh.Close.shift(14)
8.
9. In [16]: ma14 = ma.smaCal(zgyh.Close, 14)
10.
11. In [17]: data = pd.concat([zgyh.Close, mom14, ma14], 1).dropna()
12.
13. In [18]: data.columns = ['Close', 'Mom', 'SMA']
14.
15. In [19]: axis1 = plt.subplot()
16.     ...:
17.     ...: plot1, = axis1.plot(data.index, data.Close, '-r', label='Close')
18.     ...: plot2, = axis1.plot(data.index, data.SMA, '-g', label='SMA')
19.     ...: axis2 = axis1.twinx()
20.     ...:
21.     ...: plot3, = axis2.plot(data.index, data.Mom, '-b', label='Mom')
22.     ...:
23.     ...: plt.legend(handles=[plot1, plot2, plot3])

```

```
20 Out[19]: <matplotlib.legend.Legend at 0x7f25cc421240>
```

```
1 In [20]: ret = (zgyh.Close-zgyh.Close.shift(1))/zgyh.Close.shift(1)
2
3 In [21]: signal = np.all([data.Close > data.SMA, data.Mom>0],0)
4
5 In [22]: signal = pd.Series(signal,index = ret[14:].index)
6
7 In [23]: ret = ret[15:] * signal.shift(1)[1:]
8
9 In [24]: ret.sum()
10 Out[24]: -0.2389164532938523
```

当价格处于均线上方时，均线可以构成对价格的有效支撑，而动量大于 0 则表明股价正在上升。因此，我们可以设计这样的一个交易策略：当价格在均线上方且动量大于 0 时，买入。当然，这只是一个简单的想法，指标和参数的选择都较为随意，所以最后的收益还是负的。

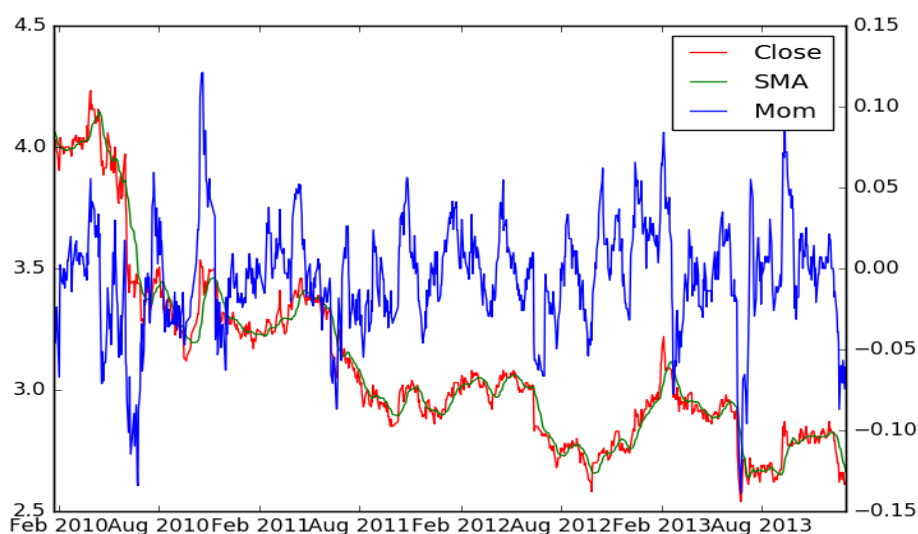


图 30.2: 题 2

(b)

```
1 In [1]: aapl = pd.read_csv('problem30-3.csv', index_col='date')
2
3 In [2]: aapl.index.name = 'Date'
4
5 In [3]: aapl.index = pd.to_datetime(aapl.index, format='%Y-%m-%d')
```

(a)

```
1 In [4]: import candle
2
```

```

3 In [5]: candle.candleLinePlots(aapl[ 2014-01 : 2014-02 ],splitFigures=False,
4 candleTitle= 'Candle Plot of AAPL',Data=aapl.Close[ 2014-01 : 2014-02 ],)

```

```

1 #因为苹果发放了股利，为了价格的一致性，我们选取了Adjusted Price，即复权过的价
  格
2 In [6]: dif = ma.ewmaCal(aapl1.Adjusted,12,2/13)-ma.ewmaCal(aapl1.Adjusted
  ,26,2/27)
3
4 In [7]: dif = dif[25:]
5
6 In [8]: dea = ma.ewmaCal(dif,9,2/10)
7
8 In [9]: dea = dea[8:]
9
10 In [10]: dif = dif[8:]
11
12 In [11]: macd = dif -dea
13
14 In [12]: plt.subplot(211)
15     ...: plt.plot(dif[12:],label=' DIF',color=' k ')
16     ...: plt.plot(dea[12:], label=' DEA',color=' b',linestyle=' dashed ')
17     ...: plt.title(" DIF and DEA ")
18     ...: plt.legend()
19     ...: plt.subplot(212)
20     ...: plt.bar(left=macd[12:].index,height=macd[12:],
21     ...:          label=' MACD',color=' r ')
22     ...: plt.legend()
23 Out[12]: <matplotlib.legend.Legend at 0x7f25cc6d9ac8>

```

(b)

```

1 In [13]: price = aapl1.Adjusted
2
3 In [14]: dat = pd.concat([price,dea,dif],1).dropna()
4
5 In [15]: dat.columns = [ 'price' , 'dea' , 'dif' ]
6
7 In [16]: asset = 100000 * np.ones(len(dat))
8
9 In [17]: cash = 100000 * np.ones(len(dat))
10
11 In [18]: share = np.zeros(len(dat))
12

```

```

13 In [19]: for i in range(1,len(dat)):
14     ...:     cash[i]=cash[i-1]+share[i-1]*dat.price[i]
15     ...:     asset[i]=cash[i]
16     ...:     if dat.dif[i] > dat.dea[i] >0 and dat.dif[i-1] < dat.dea[i-1]:
17     ...:         share[i]=100
18     ...:         cash[i]=cash[i]-100*dat.price[i]
19     ...:     elif dat.dif[i] < dat.dea[i] < 0 and dat.dif[i-1] > dat.dea[i-1]:
20     ...:         share[i]=-100
21     ...:         cash[i]=cash[i]+100*dat.price[i]
22     ...:
23     ...:
24
25 In [20]: asset[-5:]
26 Out[20]: array([ 100022.3615,  100183.0869,  100183.0869,  100183.0869,
                  100183.0869])

```

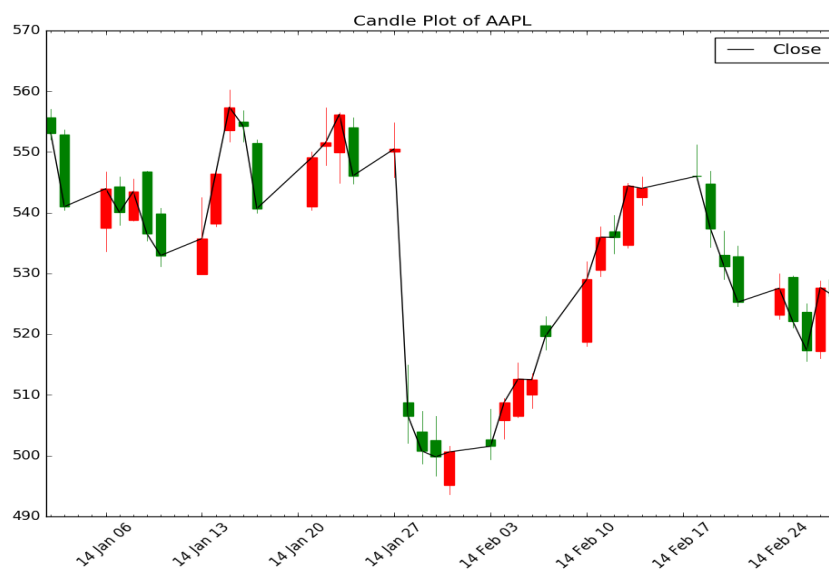


图 30.3: 题 3

4.

```

(a)
1 In [21]: SMA6 = ma.smaCal(aapl.Close,6)
2
3 In [22]: BIAS6 = (aapl.Close-SMA6)/SMA6*100

```

```

1 In [23]: EMA9 = ma.ewmaCal(aapl.Close,9)
2

```

```
3 In [24]: SMA38 = ma.smaCal(aapl.Close,38)
4
5 In [25]: BIAS = (EMA9 - SMA38)/SMA38
```

第三十一章 通道突破

(b) (a)

```
1 In [1]: import pandas as pd
2     ...:
3     ...: import numpy as np
4     ...:
5     ...: import matplotlib.pyplot as plt
6     ...:
7     ...: import candle
8     ...: def bbands(tsPrice, period=20, times=2):
9     ...:     upBBand=pd.Series(0.0, index=tsPrice.index)
10    ...:     midBBand=pd.Series(0.0, index=tsPrice.index)
11    ...:     downBBand=pd.Series(0.0, index=tsPrice.index)
12    ...:     sigma=pd.Series(0.0, index=tsPrice.index)
13    ...:     for i in range(period-1, len(tsPrice)):
14    ...:         midBBand[i]=np.nanmean(tsPrice[i-(period-1):(i+1)])
15    ...:         sigma[i]=np.nanstd(tsPrice[i-(period-1):(i+1)])
16    ...:         upBBand[i]=midBBand[i]+times*sigma[i]
17    ...:         downBBand[i]=midBBand[i]-times*sigma[i]
18    ...:     BBands=pd.DataFrame({' upBBand': upBBand[(period-1):], \
19    ...:                           ' midBBand': midBBand[(period-1):], \
20    ...:                           ' downBBand': downBBand[(period-1):], \
21    ...:                           ' sigma': sigma[(period-1):]})
22    ...:     return(BBands)
23    ...:
24    ...: Bands)
25    ...:
26
27 In [2]: GSPC = pd.read_csv('Data/Part5/005/problem31-1.csv', index_col='date'
28 )
29
30 In [3]: GSPC.index.name='Date'
31
32 In [4]: GSPC.index = pd.to_datetime(GSPC.index, format='%Y-%m-%d')
```

```

32
33 In [5]: GSPC1 = GSPC['2014-02']
34
35 In [6]: candle.candleLinePlots(GSPC1, candleTitle= 'Candle Plot of SPY500', Data
    =GSPC1.Close)

1 In [7]: candle.candleLinePlots(GSPC1, candleTitle= 'Candle Plot of SPY500',
2     ...: Data=bbands(GSPC1.Close) [['downBBand', 'upBBand']])

```

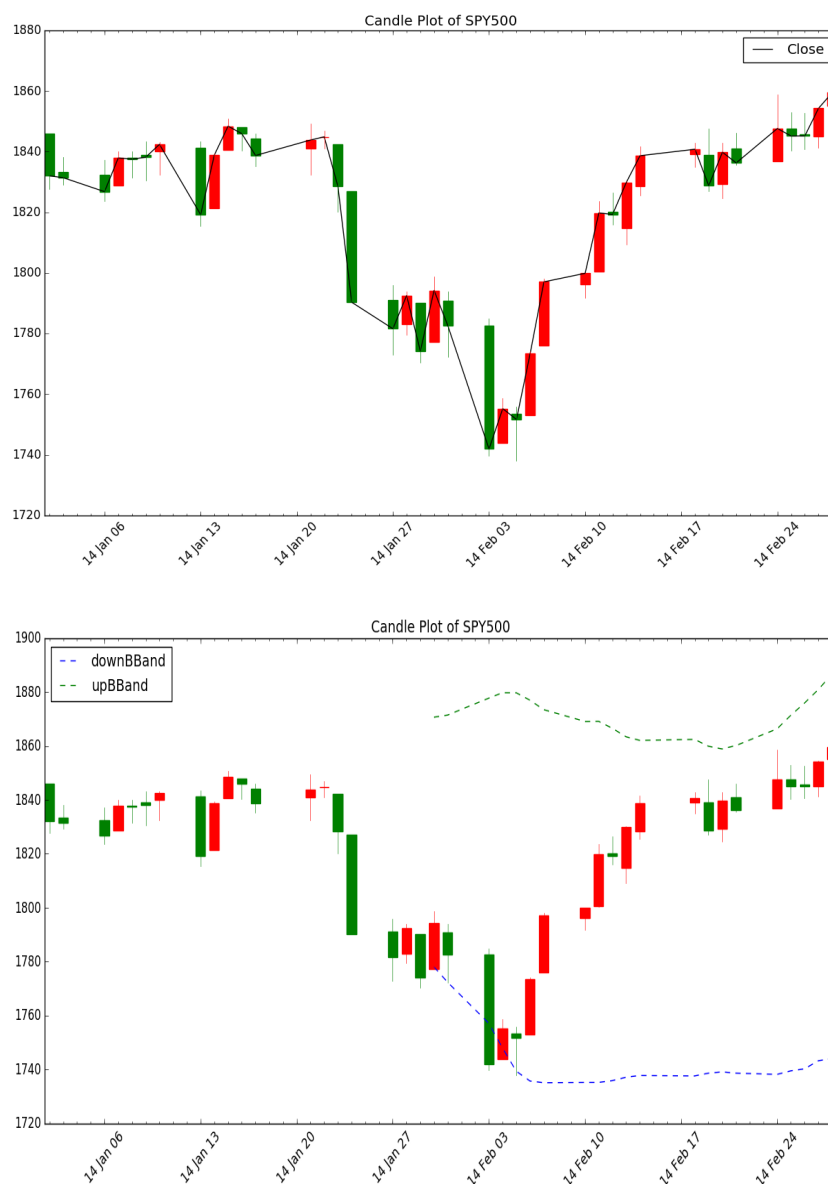


图 31.1: 题 1

收盘价与布林带中轨线具有大致相同的趋势，中轨线会稍微更加平滑一些。而上下轨线之间的距离则受到最近几天收盘价波动幅度的影响。最近几天收盘价的波动幅度越大，上下轨线之间的距离也就越大。

```

(b)
1 In [8]: upbound = pd.rolling_max(GSPC.High,10)

```

```
2
3 In [9]: lowbound = pd.rolling_min(GSPC.Low,10)
4
5 In [10]: midbound = (upbound+lowbound)/2
6
7 In [11]: bounds = pd.concat([upbound,lowbound,midbound],1)
8
9 In [12]: bounds.columns=[ upbound ; lowbound ; midbound ]
10
11 In [13]: bounds.dropna().plot()
12 Out[13]: <matplotlib.axes._subplots.AxesSubplot at 0x7f8465df1550>
```

```
1 In [14]: std = pd.rolling_std(GSPC.Close,10)
2
3 In [15]: midbound = pd.rolling_mean(GSPC.Close,10)
4
5 In [16]: upbound = midbound + 1.5 * std
6
7 In [17]: lowbound = midbound - 1.5 * std
8
9 In [18]: bounds = pd.concat([upbound,lowbound,midbound],1)
10
11 In [19]: bounds.columns=[ upbound ; lowbound ; midbound ]
12
13 In [20]: bounds.dropna().plot()
14 Out[20]: <matplotlib.axes._subplots.AxesSubplot at 0x7f8465d9dc50>
```

(d) 两者的原理基本相同，唯一的不同在于所使用的价格数据。布林带使用的是调整过后的价格，而 (d) 问使用的仅仅是收盘价。

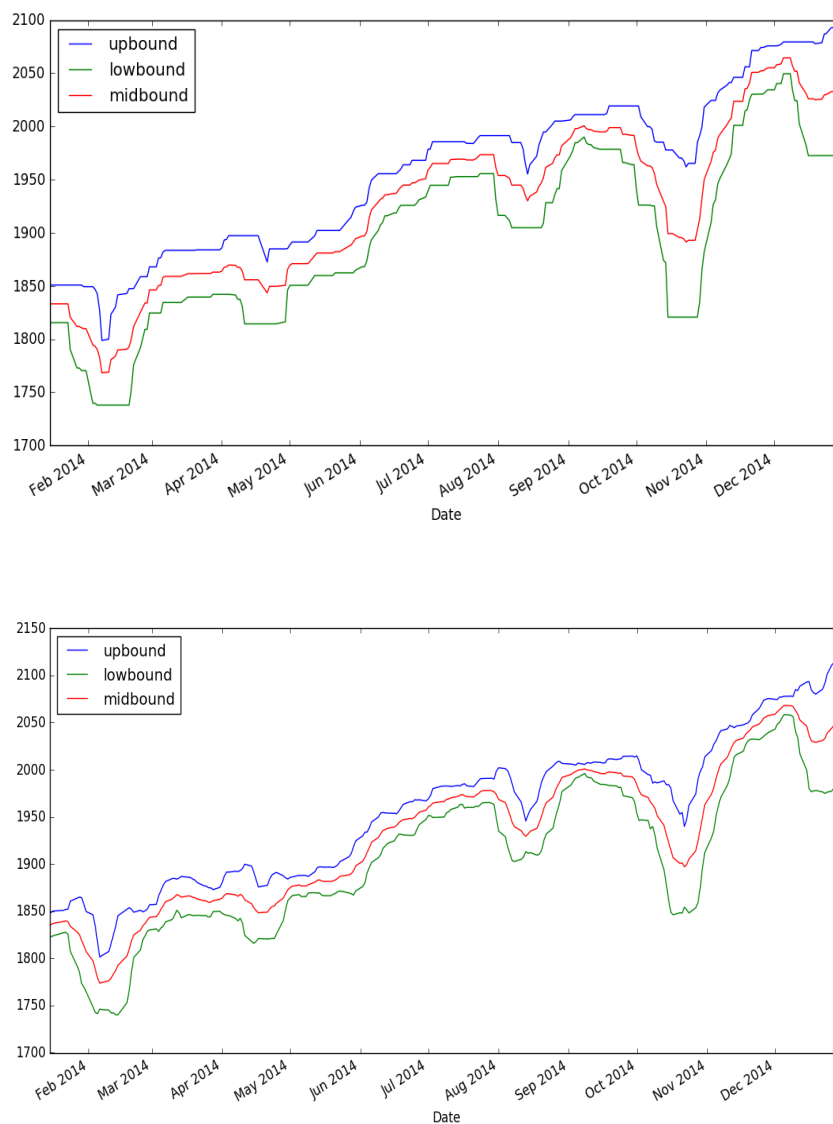


图 31.2: 题 1

2.

```

(a)
1 In [21]: bounds = bbands(GSPC.Close)
2
3 In [22]: bvalue = (GSPC.Close - bounds.downBBand) / (bounds.upBBand - bounds.
4           downBBand)
5
6 In [23]: bvalue.plot()
7 Out[23]: <matplotlib.axes._subplots.AxesSubplot at 0x7f84644af080>
8
9
10
11 In [24]: bvalue = bvalue.dropna()
12
13 In [25]: dat = pd.concat([GSPC.Close, bvalue], 1).dropna()

```

```

4
5 In [26]: dat.columns = [ 'bvalue' , 'price' ]
6
7 In [27]: asset = 100000 * np.ones(len(dat))
8
9 In [28]: cash = 100000 * np.ones(len(dat))
10
11 In [29]: share = np.zeros(len(dat))
12
13 In [30]: for i in range(1,len(dat)):
14     ...:     cash[i]=cash[i-1]+share[i-1]*dat.price[i]
15     ...:     asset[i]=cash[i]
16     ...:     if dat.bvalue[i-1]<0:
17     ...:         share[i]=100
18     ...:         cash[i]=cash[i]-100*dat.price[i]
19     ...:     elif dat.bvalue[i-1]>0:
20     ...:         share[i]=-100
21     ...:         cash[i]=cash[i]+100*dat.price[i]
22     ...:
23     ...:

```

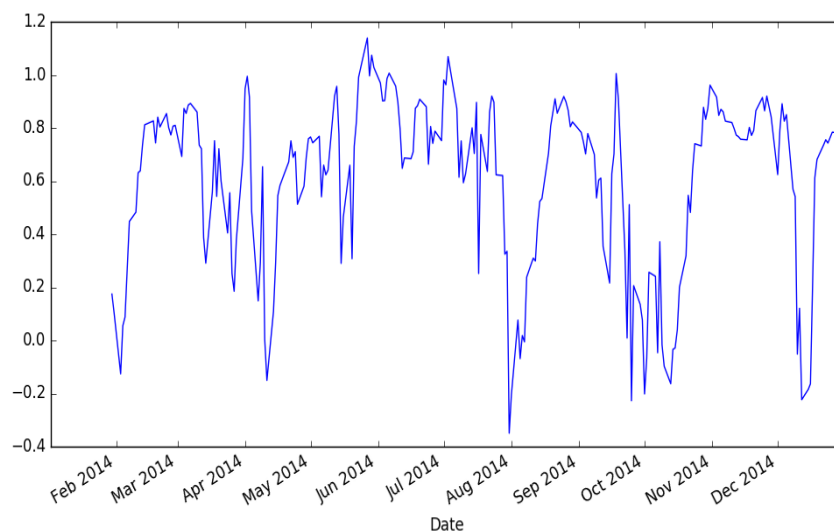


图 31.3: 题 2

(B)

(a)

```

1 In [31]: BW = (bounds.upBBand - bounds.downBBand)/bounds.midBBand

```

```

1 In [32]: BW.describe()
2 Out[32]:
3 count      233.000000
4 mean        0.046497
5 std         0.023566
6 min         0.012923
7 25%         0.028057
8 50%         0.042070
9 75%         0.061352
10 max        0.119867
11 dtype: float64

```

(1a) 从布林带的走势上看，标普 500 指数整体上是呈现出上涨的趋势。也就是说，长期上，标普 500 指数仍会上涨。从 BW 的分布上看，在大部分的时间中，BW 值较小，市场较为平稳。但是，BW 的最大值表明仍然存在着一些极端值。结合布林带的图像，我们可以发现这些极端值主要由 10 月份之后的振荡造成。综合来说，长期上，标普 500 指数仍处于上涨趋势中。但是，短期上，标普 500 指数正处于由距离振荡逐渐恢复平稳的过程中，存在着较高的风险。

```

1 In [33]: midbound = (pd.rolling_mean(GSPC.Close,3) + pd.rolling_mean(GSPC.Close
2           ...,          + pd.rolling_mean(GSPC.Close,12) + pd.rolling_mean(GSPC.Close
3           ,24))/4
4 In [34]: std = pd.rolling_std(midbound,11)
5
6 In [35]: upbound = midbound + 6 * std
7
8 In [36]: lowbound = midbound - 6 * std
9
10 In [37]: bounds = pd.concat([upbound,lowbound,midbound],1)
11
12 In [38]: bounds.columns=[ upbound , lowbound , midbound ]
13
14 In [39]: bounds.dropna().plot()
15 Out[39]: <matplotlib.axes._subplots.AxesSubplot at 0x7f8464405860>

```

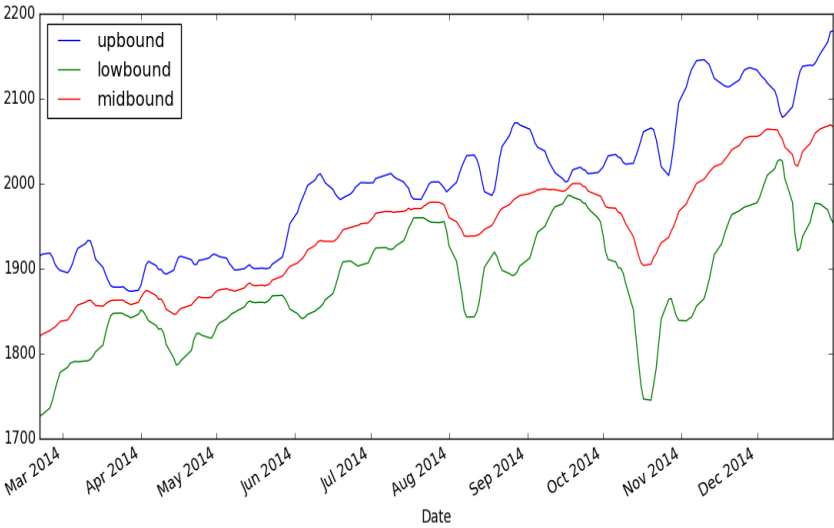


图 31.4: 题 4

第三十二章 KDJ

4. 略

2.

```
1 In [1]: import pandas as pd
2     ...:
3     ...: import numpy as np
4     ...:
5     ...: import movingAverage as ma
6     ...:
7     ...: import matplotlib.pyplot as plt
8
9 In [2]: GSPC = pd.read_csv(' Data/Part5/006/problem32-2.csv', index_col= 'date' )
10
11 In [3]: GSPC.index.name= 'Date'
12
13 In [4]: GSPC.index = pd.to_datetime(GSPC.index, format= '%Y-%m-%d' )
14
15 In [5]: min_value = GSPC.Low.rolling(9).min()
16
17 In [6]: max_value = GSPC.High.rolling(9).max()
18
19 In [7]: RSV = (GSPC.Close - min_value)/(max_value - min_value) * 100
20
21 In [8]: RSV = RSV.dropna()
22
23 In [9]: K = ma.ewmaCal(RSV,2,1/3)[1:]
24
25 In [10]: D = ma.ewmaCal(K,2,1/3)[1:]
26
27 In [11]: J = 3*K - 2*D
28
29 In [12]: J = J.dropna()
30
31 In [13]: KDJ = pd.concat([K,D,J],1).dropna()
```

```

32
33 In [14]: KDJ.columns = [ 'K' , 'D' , 'J' ]
34
35 In [15]: signal1 = [1 if KDJ.K[i] < 20 else -1 if KDJ.K[i] > 80 else 0 for i in
    range(len(KDJ))]
36
37 In [16]: signal2 = [1 if KDJ.D[i] < 20 else -1 if KDJ.D[i] > 80 else 0 for i in
    range(len(KDJ))]
38
39 In [17]: signal3 = [1 if KDJ.J[i] < 0 else -1 if KDJ.J[i] > 100 else 0 for i in
    range(len(KDJ))]
40
41 In [18]: signal = np.array(signal1) + np.array(signal2) + np.array(signal3)
42
43 In [19]: signal = pd.Series(signal, index = KDJ.index)
44
45 In [20]: ret = (GSPC.Close - GSPC.Close.shift(1))/GSPC.Close.shift(1)
46
47 In [21]: KDJRet = ret * signal.shift(1)
48
49 In [22]: KDJRet = KDJRet.dropna()
50
51 In [23]: sum(KDJRet>0)/sum(KDJRet!=0)
52 Out[23]: 0.45977011494252873

```

3.

```

(a)
1 In [24]: minValue = GSPC.Low.rolling(14).min()
2
3 In [25]: maxValue = GSPC.High.rolling(14).max()
4
5 In [26]: WR = (maxValue - GSPC.Close) / (maxValue - minValue) * 100
6
7 In [27]: WR = WR.dropna()

1 In [28]: signal = [1 if WR[i]>80 else -1 if WR[i]<20 else 0 for i in range(
    len(WR))]
2
3 In [29]: signal = pd.Series(signal, index = WR.index)
4
5 In [30]: WRRet = ret * signal.shift(1)

```

```
6
7 In [31]: WRRet = WRRet.dropna()
8
9 In [32]: sum(WRRet>0)/sum(signal!=0)
10 Out[32]: 0.48148148148148145
```

第三十三章 量价关系

(b)

```
1 In [1]: import pandas as pd
2         ...:
3         ...: import numpy as np
4         ...:
5         ...: import movingAverage as ma
6         ...:
7         ...: import matplotlib.pyplot as plt
8         ...:
9         ...: import candle
10
11 In [2]: cjzq = pd.read_csv(' Data/Part5/007/problem33-1.csv', index_col= 'date' )
12
13 In [3]: cjzq.index.name= 'Date'
14
15 In [4]: cjzq.index = pd.to_datetime(cjzq.index, format= '%Y-%m-%d' )
16
17 In [5]: candle.candleVolume(cjzq, candletitle= 'Candle Plot of Changjiang
        Securities', bartitle= 'Volume of Changjiang Securities' )
```

比如，1月27号到1月30号出现价涨量增的现象，而之后，股价也连续多日上涨。

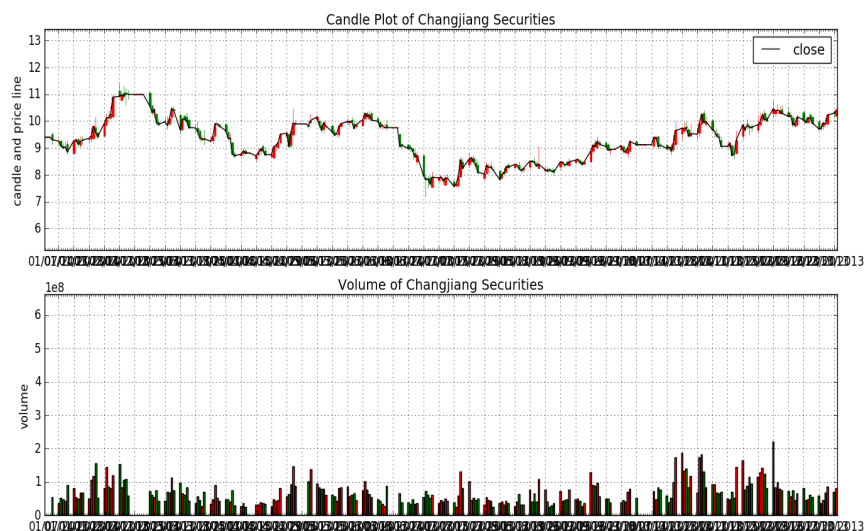


图 33.1: 题 1

```

2.
1 In [6]: close = cjzq.Close
2
3 In [7]: close.describe()
4 Out[7]:
5 count      260.000000
6 mean        9.335538
7 std         0.784656
8 min         7.570000
9 25%         8.802500
10 50%         9.385000
11 75%         9.900000
12 max         11.040000
13 Name: Close, dtype: float64
14
15 In [8]: BreakClose = np.ceil(close)
16
17 In [9]: BreakClose.name = 'BreakClose'
18
19 In [10]: volume = cjzq.Volume
20
21 In [11]: PrcChange = close.diff()
22
23 In [12]: UpVol = volume.replace(volume[PrcChange>0],0)
24
25 In [13]: UpVol[0] = 0

```

```
26
27 In [14]: DownVol = volume.replace(volume[PrcChange<=0],0)
28
29 In [15]: DownVol[0] = 0
30
31 In [16]: def VOblock(vol):
32     ...:     return([np.sum(vol[BreakClose==x]) for x in range(6,12)])
33     ...:
34     ...:
35
36 In [17]: cumUpVol = VOblock(UpVol)
37
38 In [18]: cumDownVol = VOblock(DownVol)
39
40 In [19]: ALLVol = np.array([cumUpVol,cumDownVol]).transpose()
41
42 In [20]: fig,ax=plt.subplots()
43     ...:
44     ...: ax1=ax.twinx()
45     ...:
46     ...: ax.plot(close)
47     ...:
48     ...: ax.set_title(' Cumulative volume in different price intervals' )
49     ...:
50     ...: ax.set_ylim(6,12)
51     ...:
52     ...: ax.set_xlabel(' time' )
53     ...:
54     ...: plt.setp(ax.get_xticklabels(), rotation=20, horizontalalignment=' center' )
55     ...:
56     ...: ax1.barh(bottom=range(6,12),width=ALLVol[:,0],
57     ...:           height=1,color=' g ',alpha=0.2)
58     ...:
59     ...: ax1.barh(bottom=range(6,12),width=ALLVol[:,1],height=1,left=ALLVol[:,0],
60     ...:           color=' r ',alpha=0.2)
61 Out[20]: <Container object of 6 artists>
```

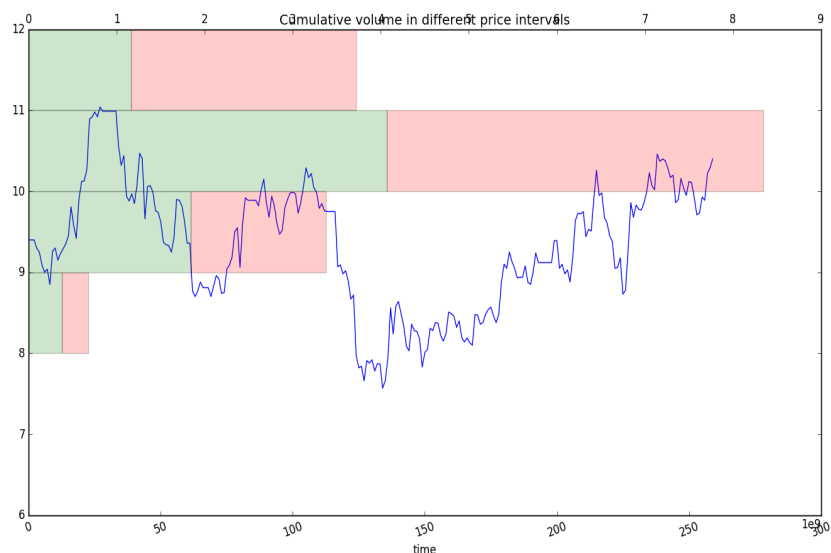


图 33.2: 题 2

```

3. In [21]: VoMA = np.zeros(len(cjzq))
1
2
3 In [22]: for i in range(5, len(cjzq)):
4     ...:     Volume = volume[i-5:i]
5     ...:     Close = close[i-5:i]
6     ...:     VoMA[i] = sum(Close * Volume / sum(Volume))
7     ...:
8     ...:
9
10 In [23]: VoMA = pd.Series(VoMA, index = cjzq.index).replace(0, np.nan).dropna()
11
12 In [24]: VoMA.plot()
13 Out[24]: <matplotlib.axes._subplots.AxesSubplot at 0x7f16bc5ec208>

```

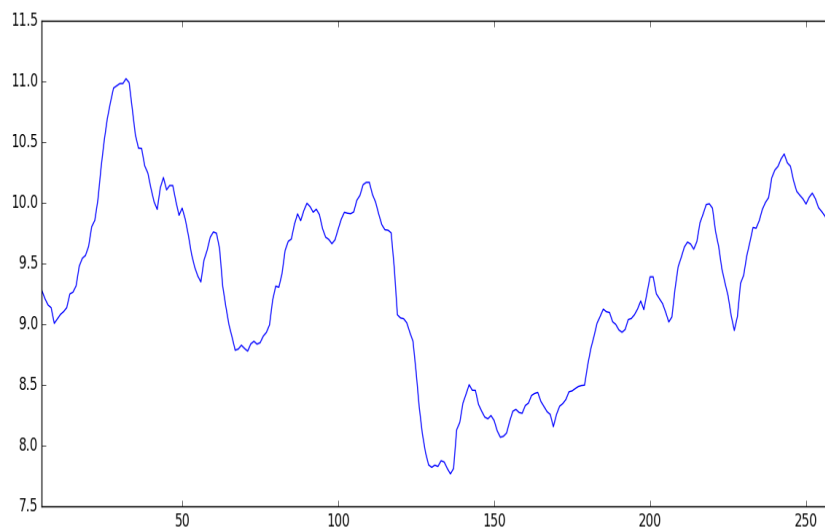


图 33.3: 题 3

```

4.
1 In [25]: midbound = VoMA
2
3 In [26]: upbound = midbound + 1.5 * pd.rolling_std(midbound,6)
4
5 In [27]: lowbound = midbound - 1.5 * pd.rolling_std(midbound,6)
6
7 In [28]: dat = pd.concat([close,upbound,lowbound],1).dropna()
8
9 In [29]: dat.columns = [ ' price' , upbound , lowbound ]
10
11 In [30]: asset = 2000 * np.ones(len(dat))
12
13 In [31]: cash = 2000 * np.ones(len(dat))
14
15 In [32]: share = np.zeros(len(dat))
16
17 In [33]: for i in range(1,len(dat)):
18     ...:     cash[i]=cash[i-1]+share[i-1]*dat.price[i]
19     ...:     asset[i]=cash[i]
20     ...:     if dat.price[i-1]< dat.upbound[i-1]:
21     ...:         share[i]=100
22     ...:         cash[i]=cash[i]-100*dat.price[i]
23     ...:     elif dat.price[i-1]>dat.lowbound[i-1]:
24     ...:         share[i]=-100
25     ...:         cash[i]=cash[i]+100*dat.price[i]

```

26	...
27	...

第三十四章 OBV

```
1. In [1]: import pandas as pd
2
3 In [2]: import numpy as np
4
5 In [3]: she = pd.read_csv('Data/Part5/008/problem34-1.csv' ,
6 ...:                      index_col='date' )
7
8 In [4]: she.index.name='Date'
9
10 In [5]: she.index = pd.to_datetime(she.index, format='%Y-%m-%d' )
11
12 In [6]: MendOBV = (she.Volume * ((she.Close-she.Low) - (she.High-she.Close))\
13 ...:              /(she.High - she.Low)).cumsum().dropna()
14
15 In [7]: signal = pd.Series(np.where(MendOBV.diff()>0 , 1, -1),index = MendOBV.
16 ...:                       index)[1:]
17
18 In [8]: ret = (she.Close - she.Close.shift(1))/she.Close.shift(1)
19
20 In [9]: OBVret = ret[2:] * signal.shift(1)[1:]
21 #计算获胜率
22 In [10]: sum(OBVret>0)/len(signal)
23 Out[10]: 0.48589341692789967
24
25 #模拟资产账户
26 In [11]: %paste
27 dat = pd.concat([signal,she.Close],1).dropna()
28 dat.columns = ['signal' , 'price' ]
29
30 asset = 10000 * np.ones(len(dat))
31 cash = 10000 * np.ones(len(dat))
32 share = np.zeros(len(dat))
```

```

32 for i in range(1,len(dat)):
33     cash[i]=cash[i-1]+share[i-1]*dat.price[i]
34     asset[i]=cash[i]
35     if dat.signal[i-1] == 1:
36         share[i]=100
37         cash[i]=cash[i]-100*dat.price[i]
38     elif dat.signal[i-1] == -1:
39         share[i]=-100
40         cash[i]=cash[i]+100*dat.price[i]
41
42 ## --- End pasted text ---

```

```

2.
1 In [12]: %paste
2 difClose = she.Close.diff()
3
4 difClose[0] = 0
5
6 volume = she.Volume
7
8 volume = volume.replace(0, np.nan).dropna()
9
10 difClose = difClose[volume.index]
11
12 OBV=(((difClose>=0)*2-1) * she.Volume).cumsum()
13
14 dat = pd.concat([OBV,she.Close],1).dropna()
15
16 dat.columns = [ 'OBV' , 'price' ]
17
18 asset = 10000 * np.ones(len(dat))
19
20 cash = 10000 * np.ones(len(dat))
21
22 share = np.zeros(len(dat))
23
24 for i in range(2,len(dat)):
25     cash[i]=cash[i-1]+share[i-1]*dat.price[i]
26     asset[i]=cash[i]
27     if dat.OBV[i-2] < 0 and dat.OBV[i-1] > 0:
28         share[i]=100

```

```

29         cash[i]=cash[i]-100*dat.price[i]
30     elif dat.OBV[i-2] > 0 and dat.OBV[i-1] < 0:
31         share[i]=-100
32         cash[i]=cash[i]+100*dat.price[i]
33
34 ## — End pasted text —

```

```

3. In [13]: %paste
1  wanke = pd.read_csv(' Data/Part5/008/problem34-3.csv ',
2                      index_col= date )
3
4
5  wanke.index.name= Date
6
7  wanke.index = pd.to_datetime(wanke.index,format= %Y-%m-%d )
8
9  difClose = wanke.Close.diff()
10
11 difClose[0] = 0
12
13 volume = wanke.Volume
14
15 volume = volume.replace(0, np.nan).dropna()
16
17 difClose = difClose[volume.index]
18
19 OBV=(((difClose>=0)*2-1) * wanke.Volume).cumsum()
20
21 smOBV = pd.rolling_mean(OBV,12)
22
23 ## — End pasted text —

```