



# Pandas基础课程

# 目录

---



# Pandas 简介

Pandas的名称来自于面板数据（panel data）和python数据分析（data analysis）。

Pandas 最初由AQR Capital Management于2008年4月开发，并于2009年底开源出来，主要提供高性能易用数据类型和分析工具

Pandas的优点有哪些？

- 可以轻易的处理浮点及非浮点数据类型的缺失值(NaN)
- 基于智能标签的切片，花式索引，轻易从大数据集中取出子集
- 可以灵活处理时间数据
- 灵活强大的分组功能，可对数据集进行拆分组合操作
- 方便的将其他Python和Pandas数据结构中不同类索引的数据转换为DataFrame对象
- 轴(axes)的分层标签（使每个元组有多个标签成为可能）



*AQR Capital Management*

创始人 *Cliff Asnester*

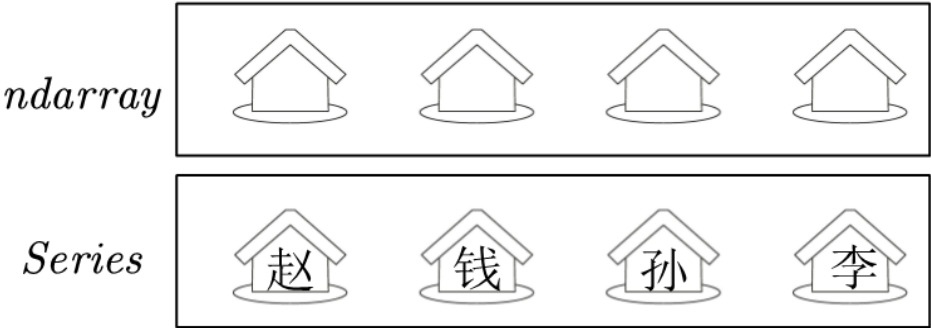
# Pandas预备知识

- Pandas一般默认简写: pd
- Pandas有两种主要的数据类型: Series和DataFrame
  - Series可以理解成一种具有索引的数组
  - DataFrame是由共用相同索引的一组列组成的
- Pandas同样适用广播机制

|   | A | B | C |
|---|---|---|---|
| 1 | 1 | 2 | 3 |
| 2 | 4 | 5 | 6 |
| 3 | 7 | 8 | 9 |

部分excel表

|   | 这        | 是        | 列        | 索        | 引        |
|---|----------|----------|----------|----------|----------|
| 这 | 0.612501 | 0.592796 | 0.773207 | 0.504646 | 0.677181 |
| 是 | 0.756287 | 0.578716 | 0.651797 | 0.804316 | NaN      |
| 行 | 0.310771 | 0.070035 | 0.050099 | 0.812016 | 0.298376 |
| 索 | 0.319321 | 0.600389 | 0.013632 | 0.441274 | 0.175991 |
| 引 | 0.279503 | 0.565244 | 0.154203 | 0.523429 | 0.852456 |



ndarray与Series的区别

DateFrame数据示例

# Series函数相关知识

---

**Series函数** 同样适用ndarray的操作，但是输出结果仍是Series类型

**pandas.Series**(self, data=None, index=None, dtype=None, name=None, copy=False, fastpath=False)

*Data*: Series不改变数据的类型

*Index*: Series可以自定义索引，索引关键词可以使用符号或汉字

## 创建Series的几种方法

Python列表

标量值（必须有index）

Python字典

ndarray

其他

# DateFrame函数练习

---

## DataFrame函数

**pandas.DataFrame**(*data=None, index=None, columns=None, dtype=None, copy=False*)

index: index定义行索引

columns: columns定义列索引

## 创建DataFrame的几种方法

二维ndarray

由一维ndarray、列表、字典、元组、或者Series构成的字典

Series类型

其他DataFrame类型

# DataFrame查看功能

| 函数                   | 描述               |
|----------------------|------------------|
| DataFrame()          | 创建一个DataFrame对象  |
| DataFrame.values     | 返回ndarray类型的对象   |
| DataFrame.index      | 获取行索引            |
| DataFrame.columns    | 获取列索引            |
| DataFrame.axes       | 获取行及列索引          |
| DataFrame.T          | 行与列对调            |
| DataFrame.info()     | 打印DataFrame对象的信息 |
| DataFrame.head(i)    | 显示前 i 行数据        |
| DataFrame.tail(i)    | 显示后 i 行数据        |
| DataFrame.describe() | 查看数据按列的统计信息      |

# Series和DataFrame索引操作

pandas数据结构的索引与ndarray的索引类似，但是pandas数据的索引不仅仅可以是数字，还可以是字符串。

与数组不同的是，Series单个索引提取值、多个索引提取Series数据，且索引提取左闭右闭。

pandas的索引可以支持更多的功能，每个索引都是一个不可变的集合，它让数据的传递更加的安全，索引支持一些集合的方法和属性，可以辅助解决常见问题。

| 方法                            | 描述                    |
|-------------------------------|-----------------------|
| in                            | 判断索引是否存在              |
| append                        | 将额外的对象粘贴到原索引后，产生一个新索引 |
| difference\union\intersection | 计算两个索引的差集\并集\交集       |
| isin                          | 计算表示每一个值是否在传值容器的布尔数组  |
| delete                        | 将位置i的元素删除，并产生新的索引     |
| drop                          | 根据传入参数删除指定索引，并产生新的索引  |
| insert                        | 将位置i插入元素，并产生新的索引      |
| unique                        | 计算索引的唯一值序列            |

# Series和DataFrame基础操作（一）

---

Pandas需要通过改变索引进行增加和删除内部数据

添加列数据可以直接进行insert函数进行添加

**Dataframe.insert**(*loc, column, value, allow\_duplicates=False*)

*loc*: 使用整数定义\_列数据\_插入的位置, 必须是0到columns列标签的长度

*column*: 可选字符串、数字或者object; 列标签名

添加行数据需要使用reindex函数创建一个符合新索引的新对象。

**DataFrame.reindex**(*labels=None, index=None, columns=None, axis=None, method=None, copy=True, level=None, fill\_value=nan, limit=None, tolerance=None*)

*index, columns*: 新的行列自定义索引

*fill\_value*: 重新索引中用于填充缺失位置的值

*method*: 填充方法, ffill当前值向前填充, bfill向后填充

## Series和DataFrame基础操作（二）

---

对于已经建立了的Series和DataFrame数据可以通过drop函数删除指定行列数据

**DataFrame.drop** (*labels=None, axis=0, index=None, columns=None, level=None, inplace=False, errors='raise'*)

*labels*: 就是要删除的行列的名字, 用列表给定

*axis*: 默认为0, 指删除行, 因此删除columns时要指定axis=1

*index*: 直接指定要删除的行

*columns*: 直接指定要删除的列

*inplace=False*: 默认该删除操作不改变原数据, 而是返回一个执行删除操作后的新dataframe

*inplace=True*: 则会直接在原数据上进行删除操作, 删除后无法返回

# Series和DataFrame基础操作（三）

Series和DataFrame数据直接进行计算时，数据会根据索引进行计算，对于不匹配的数据直接填充NaN。

|   | a    | b    | c    |
|---|------|------|------|
| a | 0.87 | 0.23 | 0.50 |
| b | 1.00 | 0.43 | 0.11 |

|   | b    | c    | d    |
|---|------|------|------|
| b | 0.21 | 0.57 | 0.34 |
| c | 0.06 | 0.21 | 0.74 |

|   | a   | b    | c    | d   |
|---|-----|------|------|-----|
| a | NaN | NaN  | NaN  | NaN |
| b | NaN | 0.64 | 0.68 | NaN |
| c | NaN | NaN  | NaN  | NaN |

Series和DataFrame数据也支持根据条件判断进行数据选择。

|   | a    | b    | c    |
|---|------|------|------|
| a | 0.94 | 0.65 | 0.19 |
| b | 0.33 | 0.19 | 0.48 |
| c | 0.56 | 0.01 | 0.66 |

>0.5

|   | a     | b     | c     |
|---|-------|-------|-------|
| a | True  | True  | False |
| b | False | False | False |
| c | True  | False | True  |

重新  
赋值

|   | a    | b    | c    |
|---|------|------|------|
| a | 6.00 | 6.00 | 0.19 |
| b | 0.33 | 0.19 | 0.48 |
| c | 6.00 | 0.01 | 6.00 |

# Series和DataFrame基础操作（四）

---

## 索引排序

**Series.sort\_index**(*axis=0, ascending=True, inplace=False, ...*)

*axis*: 轴直接排序。对于系列，只能为0

*ascending*: 升序与降序排序

*inplace*: 如果为True，则就地执行操作

## 值排序

**DataFrame.sort\_values**(*by='##', axis=0, ascending=True, inplace=False, ...*)

*by*: 指定列名(*axis=0*或'*index*')或索引值(*axis=1*或'*columns*')

## 返回名次

**Series.rank**(*axis=0, method='average', na\_option='keep', ascending=True*)

*method*: {'average', 'min', 'max', 'first', 'dense'}指定排名时用于破坏平级关系的method选项

# Series和DataFrame切片

---

Pandas的两种数据类型都可以使用numpy一样的数据格式进行索引，也可以使用loc和iloc进行更高级的索引与切片。

**df.loc[]**：获取具有特定标签的行(和/或列)，接受两个参数：行标和列标，当列标省略时，默认获取整行数据。两个参数都可以以字符，切片以及列表的形式传入。

**df.iloc[]**：用法和df.loc[]类似，最大的区别在于，loc是基于行列的标签进行检索，而iloc是基于位置进行检索。

标签索引与内置索引混用：

**get\_loc()**是一个索引方法，意思是“获取标签在这个索引中的位置”。注意，因为用**iloc**切片不包含它的端点，必须在这个值上加上1。

# 数据读取

## 文本文件和CSV数据读取

**pandas.read\_csv**(filepath\_or\_buffer, sep=' ', header='infer', names=None, index\_col=None, usecols=None, dtype=None, skiprows=None, skipfooter=None, nrows=None, na\_values=None, skip\_blank\_lines=True, parse\_dates=False, comment=None, encoding=None)

filepath\_or\_buffer: 文件路径

sep: 指定分隔符。如果不指定参数, 则会尝试使用逗号分隔

header: 指定行数用来作为列名, 数据开始行数。如果文件中没有列名, 则默认为0, 否则设置为None

encoding: 指定字符集类型, 通常指定为'utf-8'

**\*csv换成table即可读取txt文本文件**

## Excel文件读取

**pandas.read\_excel**(io, sheet\_name=0, header=0, names=None, index\_col=None, usecols=None, squeeze=False, dtype=None, encoding=None...)

sheet\_name: 字符串用于工作表名称, 整数用于零索引工作表位置, 字符串列表或整数列表用于请求多个工作表, 为None时获取所有工作表。

# 数据存写

## 文本文件和CSV数据写入

**DataFrame.to\_csv**(*path\_or\_buf=None, sep=',', na\_rep='', columns=None, header=True, index=True, index\_label=None, mode='w', encoding=None*)

*sep*: 接收string。代表分隔符。默认为 “,”

*index\_labels*: 接收sequence。表示索引名。默认为None

*na\_rep*: 接收string。代表缺失值。默认为 ""

*columns*: 接收list。代表写出的列名。

*encoding*: 接收特定string。代表存储文件的编码格式。

*header*: 接收boolean，代表是否将列名写出。默认：True或者None

## Excel文件写入

**DataFrame.to\_excel**(*excel\_writer, sheet\_name='Sheet1', na\_rep='', float\_format=None, columns=None, header=True, index=True, index\_label=None, ...*)

# pandas统计功能

Pandas基于NumPy，同样可以使用NumPy的函数对数据框进行描述性统计

pandas还提供了更加便利的方法来计算均值，如detail['amounts'].mean()。

pandas还提供了方法叫作describe，能够一次性得出数据框所有数值型特征的非空值数目、均值、四分位数、标准差。

| 方法名称     | 说明   | 方法名称  | 说明     |
|----------|------|-------|--------|
| min      | 最小值  | max   | 最大值    |
| mean     | 均值   | ptp   | 极差     |
| median   | 中位数  | std   | 标准差    |
| var      | 方差   | cov   | 协方差    |
| sem      | 标准误差 | mode  | 众数     |
| skew     | 样本偏度 | kurt  | 样本峰度   |
| quantile | 四分位数 | count | 非空值数目  |
| describe | 描述统计 | mad   | 平均绝对离差 |

| 函数名称       | 说明    |
|------------|-------|
| np.min/max | 最小/大值 |
| np.mean    | 均值    |
| np.median  | 中位数   |
| np.var     | 方差    |
| np.ptp     | 极差    |
| np.std     | 标准差   |
| np.cov     | 协方差   |

# 目录

---



# 处理缺失值

pandas支持多种方法处理缺失值，可以直接调用函数进行处理

| 方法      | 描述                                  |
|---------|-------------------------------------|
| dropna  | 根据每个标签的值是否缺失数据来筛选标签，并根据允许缺失的数据来确定阈值 |
| fillna  | 用某些值填充缺失的数据或使用插值方法                  |
| isnull  | 返回表明哪些值是缺失值的布尔值                     |
| notnull | isnull的反函数                          |

**DataFrame.dropna**(*self*, *axis=0*, *how='any'*, *thresh=None*, *subset=None*, *inplace=False*)

*how*: 选择方式是全为空还是部分为空

**DataFrame.fillna**(*value=None*, *method=None*, *axis=None*, *inplace=False*, *limit=None*, *\*\*kwargs*)

*value*: 它是一个用于填充空值的值

*method*: 一种用于填充重新索引的Series中的空值的方法

# 处理重复值

---

**DataFrame.drop\_duplicates**(*subset=['A','B'],keep='first,inplace=True*)

*subset*: 列名, 表示只考虑这两列, 将这两列对应值相同的行进行去重。默认为*subset=None*表示考虑所有列

*keep='first'*: 保留第一次出现的重复行, 是默认值。keep另外两个取值为“last”和False, 分别表示保留最后一次出现的重复行和去除所有重复行

*inplace=True*: 表示直接在原来的DataFrame上删除重复项, 而默认值False表示生成一个副本

# 数据离散分箱

连续数值经常需要离散化，或者分离成箱子进行分析。如果相对数据进行分组，需要借用cut函数和qcut函数。

**pandas.cut**(*x, bins, right=True, labels=None, retbins=False, precision=3, duplicates='raise'*)

**pandas.qcut**(*x, bins, labels=None, retbins=False, precision=3, duplicates='raise'*)

*x*: 被切分的类数组数据，必须是1维的（不能用DataFrame）

*bins*: bins是被切割后的区间，有3中形式：一个int型的标量、标量序列（数组）或pandas.IntervalIndex

*right*: bool型参数，默认为True，表示是否包含区间右部。比如如果bins=[1,2,3], right=True, 则区间为(1,2],(2,3]; right=False, 则区间为(1,2),(2,3)

*labels*: 给分割后的bins打标签，比如把年龄*x*分割成年龄段bins后，可以给年龄段打上诸如青年、中年的标签

*retbins*: 是否返回bins

*precision*: 数据精度

*duplicates*: 是否允许重复区间。有两种选择：raise: 不允许, drop: 允许

以上内容仅为本文档的试下载部分，为可阅读页数的一半内容。如要下载或阅读全文，请访问：<https://d.book118.com/557106104043006141>