

Welcome to BIO-210

Applied software engineering for life sciences

November 18th 2024 – Lecture 10

Prof. Alexander MATHIS

EPFL

	Date	Topic	Software version	Software releases	Grading / Feedback
0	09/09/2024	Python introduction I			
1	16/09/2024	Public holiday			
2	23/09/2024	Python introduction II			
3	30/09/2024	Git and GitHub (+installation VS Code)			
4	07/10/2024	Project introduction	v1		
5	14/10/2024	Functionify	v2	v1	
6	21/10/2024	EPFL fall break			
7	28/10/2024	Visualization and documentation	v3	v2	code review (API)
8	04/11/2024	Unit-tests, functional tests	v4	v3	
9	11/11/2024	Code refactoring	v5	v4	graded (tests)
10	18/11/2024	Profiling and code optimization	v6	v5	code review
11	25/11/2024	Object oriented programming	v7	v6	graded (speed)
12	02/12/2024	Model analysis and project report	v8	v7	code review (OO)
13	09/12/2024	Work on project			
14	16/12/2024	Wrap up		v8	graded (project)

Status of your project

You have

- a readme with information on how to install the code and what it can do
- a module for relevant functions and utilities (e.g., plotting)
- tests to assess the correctness of the code (also for refactoring)
- we also gave you some test targets, so your functions should be correct
- scripts to reproduce various results
- ...

You may wonder ...

- what is the slowest part of my code?
- how often is a particular piece of code run?
- how long does it take?
- why does my code take so long ...

Profiling

Profiling addresses these questions and helps you optimize your code - so that you can make it faster.

To get the biggest speed improvement per coding hour, you should *focus on the slowest part of your program* (in a typical use case).

Note the specific use-case, can also be user-specific and one can optimize for different target groups.

We'll consider the following tools:

- cProfile
- timeit
- line_profiler

Donald Knuth's advice

From the Art of Programming:

“The real problem is that programmers have spent far too much time worrying about efficiency in the wrong places and at the wrong times; premature optimization is the root of all evil (or at least most of it) in programming.”

Fastest code prize!

- we will assess speed of key functions (see problem set)



Note on the problem set

- you should use the different profiling tools I describe today
- you should release your fastest code as v6
- if your version has any special dependencies, make that clear in the readme
- if you develop variants that ultimately aren't the fastest you can document it in Issues (for yourself & the team); we won't grade this, but it's a good exercise to learn how to keep track of different variants.
- Typically, you would leave such variants as branches (if they are interesting enough)

P.S.: We will run your code on a machine with CPUs, not GPUs. So optimize for this setting. Note: the choice of hardware can make different algorithms more efficient!

Built-in profilers: cProfile (and profile)

cProfile is recommended, as it is the C extension with less overhead than profile

- records total run time
- records the time taken by each function
- ...

→ *This allows you to find which parts need optimization*

- recorded data can be exported, looked at with the pstats module and visualized with the snakeviz module.

cProfile syntax

```
1  import cProfile          # import the module
2  cProfile.run(statement, filename=None)
```

- As `statement` you can pass python code or a function (as a string).
- Output can be saved in filename (if passed)

A demo program

What does this program do and how should the profiling output look like?

```
1  import time
2
3  def shortbreak():
4      time.sleep(0.05)
5
6  def longbreak():
7      time.sleep(1)
8
9  def program():
10     print("Starting!")
11
12     for l in range(3):
13         shortbreak()
14
15     print("Now to the long break...")
16     longbreak()
17
18     print("Done!")
```


Adding cProfiler (when executed as script)

```
1  # Script: cProfiler_demo.py
2  import time
3
4  def shortbreak():
5      time.sleep(0.05)
6
7  def longbreak():
8      time.sleep(1)
9
10 def program():
11     print("Starting!")
12     for l in range(3):
13         shortbreak()
14     print("Now to the long break...")
15     longbreak()
16     print("Done!")
17
18 if __name__ == '__main__':
19     import cProfile
20     profiler = cProfile.Profile()           # Initializing a profiler instance
21     profiler.runcall(program)              # Run it on program()
22     profiler.print_stats()                 # Print runtime statistics
```

Running the script

```
1  alex@mac% python3 cProfiler_demo.py
2  Starting!
3  Now to the long break...
4  Done!
5          13 function calls in 1.161 seconds
6
7  Ordered by: standard name
8
9  ncalls  tottime  percall  cumtime  percall  filename:lineno(function)
10         3      0.000      0.000      0.160      0.053  cProfiler_demo.py:3(shortbreak)
11         1      0.000      0.000      1.000      1.000  cProfiler_demo.py:6(longbreak)
12         1      0.000      0.000      1.161      1.161  cProfiler_demo.py:9(program)
13         3      0.000      0.000      0.000      0.000  {built-in method builtins.print}
14         4      1.161      0.290      1.161      0.290  {built-in method time.sleep}
15         1      0.000      0.000      0.000      0.000  {method 'disable' of '_lsprof.Profiler' obj
```

What is summarized in the output?

Line #1: summary (# function calls and time)

Line #2: Ordered by: standard name (i.e. ordered by function name)

In the table:

- ``ncalls``: number of calls
- ``tottime``: total time spent in the given function (and excluding time spend in calls to sub-functions)
- ``percall``: ratio `tottime/ncalls`
- ``cumtime``: cumulative time spent in this function and all subfunctions (from invocation until exit). This figure is accurate even for recursive functions.
- ``percall``: `cumtime/primitive calls` [see later!]
- ``filename:lineno``: provides the respective data of each function

You can also store/load the stats

```
1  from cProfiler_demo import program    # importing the program from prior
2  import cProfile
3
4  # You can also save stats, here to the file summary.stats
5  cProfile.run('program()', 'summary.stats')
6
7  import pstats
8  stats = pstats.Stats('summary.stats') # Load the stats
9  stats.print_stats()                  # Print them
```

Let's store this as `cProfiler\_visualization.py`

Output of this program

```
1  alex@mac% python3 cProfiler_visualization.py
2  Starting!
3  Now to the long break...
4  Done!
5  Thu Nov 14 17:21:34 2024      summary.stats
6
7      15 function calls in 1.170 seconds
8
9  Random listing order was used
10
11  ncalls  tottime  percall  cumtime  percall  filename:lineno(function)
12      1      0.000      0.000      1.170      1.170 {built-in method builtins.exec}
13      3      0.000      0.000      0.000      0.000 {built-in method builtins.print}
14      4      1.169      0.292      1.169      0.292 {built-in method time.sleep}
15      1      0.000      0.000      0.000      0.000 {method 'disable' of '_lsprof.Profiler' obj
16      1      0.000      0.000      1.004      1.004 /Users/alex/Code/Teaching/profiling/cProfil
17      1      0.000      0.000      1.170      1.170 /Users/alex/Code/Teaching/profiling/cProfil
18      3      0.000      0.000      0.165      0.055 /Users/alex/Code/Teaching/profiling/cProfil
19      1      0.000      0.000      1.170      1.170 <string>:1(<module>)
```

You can sort the output

```
1 stats.sort_stats('calls')      # sort it by calls (see link in title for more options)
2 stats.print_stats()
```

gives:

```
1 Thu Nov 14 18:31:15 2024      summary.stats
2
3      15 function calls in 1.170 seconds
4
5      Ordered by: call count
6
7      ncalls  tottime  percall  cumtime  percall  filename:lineno(function)
8           4      1.170      0.293      1.170      0.293 {built-in method time.sleep}
9           3      0.000      0.000      0.000      0.000 {built-in method builtins.print}
10          3      0.000      0.000      0.165      0.055 /Users/alex/Code/Teaching/profiling/cProfil
11          1      0.000      0.000      1.170      1.170 {built-in method builtins.exec}
12          1      0.000      0.000      0.000      0.000 {method 'disable' of '_lsprof.Profiler' obj
13          1      0.000      0.000      1.005      1.005 /Users/alex/Code/Teaching/profiling/cProfil
14          1      0.000      0.000      1.170      1.170 /Users/alex/Code/Teaching/profiling/cProfil
15          1      0.000      0.000      1.170      1.170 <string>:1(<module>)
```

Visualization of the the cProfiler output

Snakeviz provides visualizations based on profiling output.

Snakeviz is a Python package and thus needs to be installed --> `pip install snakeviz`

The visualization is interactive, and available in the browser!

Can be used from the command line for profiling output (`summary.stats`).

```
1 snakeviz summary.stats
```

Can also be used in IPython directly, see docs!

Visualization type 1 (icicle)

```
1 snakeviz summary.stats
```

SnakeViz

Reset Root

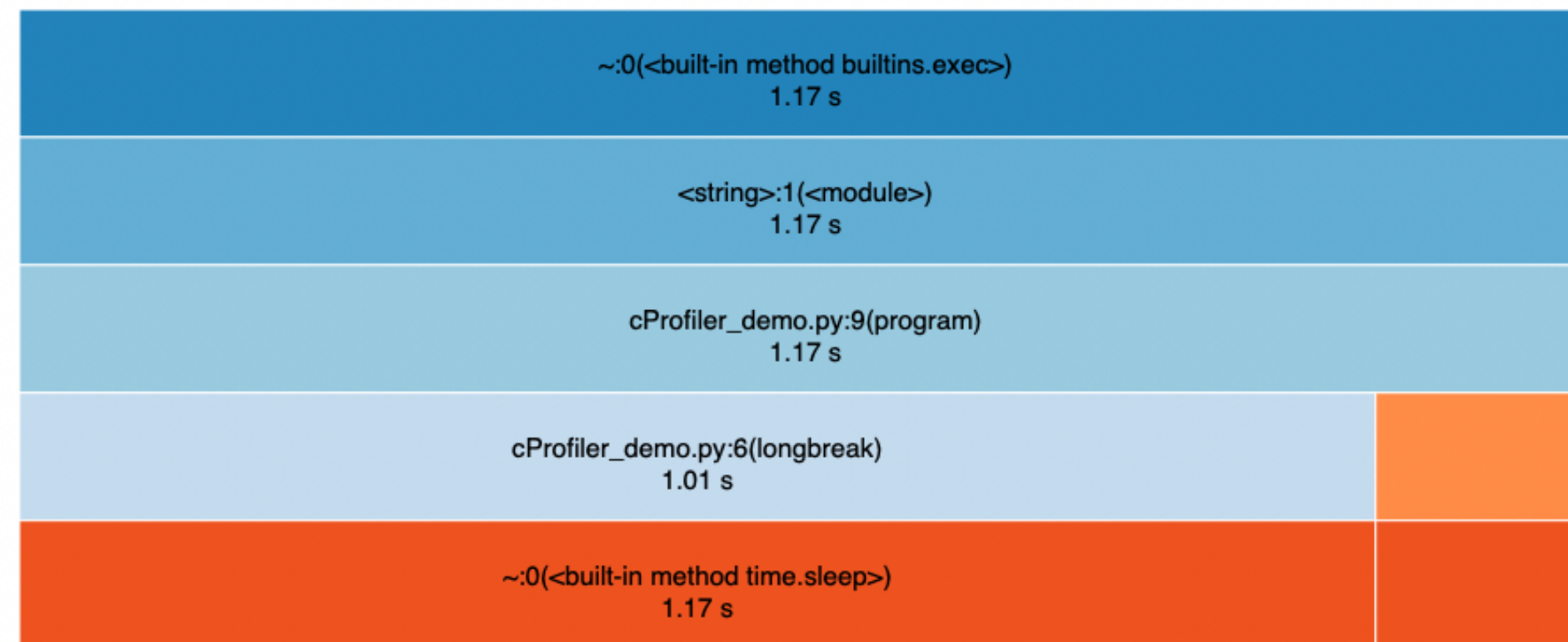
Reset Zoom

Style: **Icicle** ▾

Depth: **10** ▾

Cutoff: **1 / 1000** ▾

Call Stack



Visualization type 2 (sunburst)

1 snakeviz summary.stats

SnakeViz

Reset Root

Reset Zoom

Style: Sunburst

Depth: 15

Cutoff: 1 / 100

Name:

program

Cumulative Time:

1.17 s (100.00 %)

File:

cProfiler_demo.py

Line:

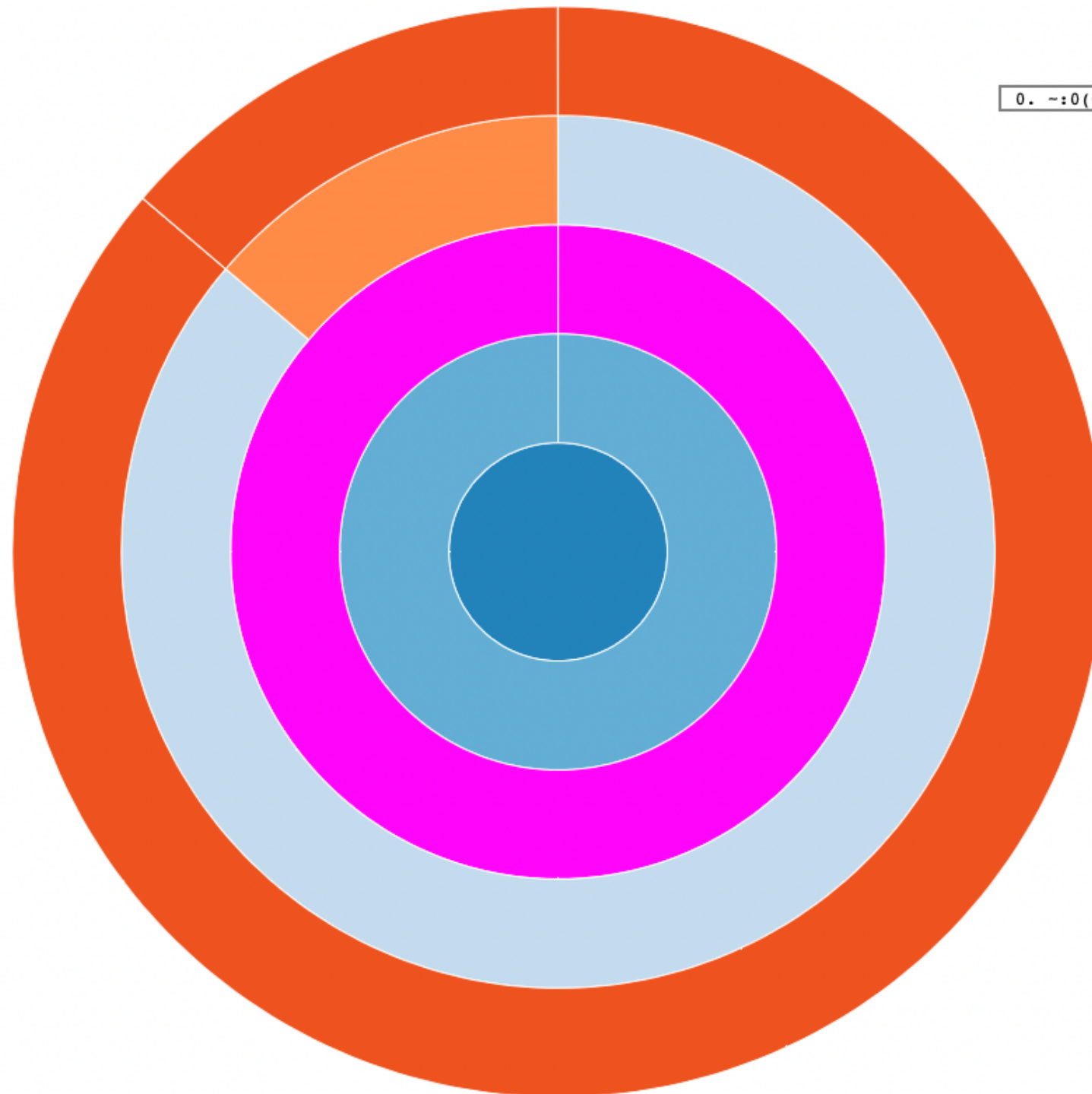
9

Directory:

/Users/alex/Code/Teaching/pr
ofiling/

Call Stack

0. ~:0(<built-in method builtins.exec>)



Alternative use of cProfile

Running cProfile on a script (in the command line).

Here sorted by total-time, and stored in summary.txt

```
1 python -m cProfile -s tottime script2profile.py > summary.txt
```

Reminder: Fibonacci

A possible solution for Fibonacci:

```
1  number_of_fibonaccielements = 50
2  # Initialization
3  x0 = 0
4  x1 = 1
5
6  fibo_list = [x0, x1]
7  i = 2
8
9  while i < number_of_fibonaccielements:
10     # Updating
11     next_x = x0 + x1          # computing sum
12     fibo_list.append(next_x) # data collection
13     x0 = x1                  # updating 1
14     x1 = next_x              # updating 2
15     i += 1
16
17  print(fibo_list) # result
```

Note: $F_0 = 0$, $F_1 = 1$, $F_n = F_{n-1} + F_{n-2}$ for $n > 1$

Python allows recursive functions

$$F_0 = 0, F_1 = 1,$$

$$F_n = F_{n-1} + F_{n-2} \text{ for } n > 1$$

You can define this as follows:

```
1  def fib(n):  
2      if n < 2:  
3          return n  
4      return fib(n-1) + fib(n-2)
```


What will ncalls be?

```
1  def fib(n):
2      if n < 2:
3          return n
4      return fib(n-1) + fib(n-2)
5
6  if __name__ == "__main__":
7      import cProfile
8      profiler = cProfile.Profile() # Initializing a profiler instance
9      profiler.run('fib(32)')       # Run profiler on fib(32)
10     profiler.print_stats()        # Print stats
```

Profiling results for `fib(32)`

```
1 7049158 function calls (4 primitive calls) in 1.148 seconds
2
3 Ordered by: standard name
4
5 ncalls      tottime  percall  cumtime  percall filename:lineno(function)
6 7049155/1    1.148    0.000    1.148    1.148 <ipython-input-14-cb3508afaaa3>:1(fib)
7 1           0.000    0.000    1.148    1.148 <string>:1(<module>)
8 1           0.000    0.000    1.148    1.148 {built-in method builtins.exec}
9 1           0.000    0.000    0.000    0.000 {method 'disable' of '_lsprof.Profiler' obj}
```

Note: **primitive call** counts calls excluding recursion.

cProfile may have a lot of overhead

when ncalls >> 1!

The recursive Fibonacci implementation nicely can illustrate the overhead costs of cProfile. Let's measure it it with `timeit.repeat`

```
1  ## adding to if __name__ == "__main__":
2  import timeit
3  print("Runtime:", timeit.repeat('fib(32)', \
4                                  setup="from __main__ import fib", repeat= 5, number=1))
```

returns (5 `repeats` of the runtime as a list; `number` is # of iterations for the main loop to be carried out):

```
1  Runtime: [0.5889114580000125, 0.588767292, 0.5895812910000018,
0.5903672500000141, 0.5895792499999857]
```

which is about half as long!!! Thus,

- don't use cProfile for estimating runtimes (per se)
- use it for assessing relative runtimes & bottlenecks

Alternatives to cProfile: Statistical profiling

- Statical profiler do not track every call, but instead the call stack every 1ms
- this gives much lower overhead costs
- note that tracking profilers (such as cProfile) can distort the results (e.g. when there are lots of lots of function calls)
- If you are interested, check out the statistical profiler Pyinstrument

Timeit is a lightweight alternative for us!

Timeit: measure execution time for code snippets

The Timeit module provides a simple way to time small bits of Python code (it is highly accurate and runs code statements multiple times for robustness).

Simple example:

```
1  # testing timeit()
2  In [1]: import timeit
3          ...: print(timeit.timeit('output = 30*125'))
4  0.032772750000003015
5  # Combining multiple lines with ;
6  In [2]: print("Summation takes ",timeit.timeit(stmt='x=123.3;y=0.245;sum=x+y'))
7  Summation takes  0.066319874999999781
```


Timeit: measure execution time for code snippets

The Timeit module provides a simple way to time small bits of Python code (it is highly accurate and runs code statements multiple times for robustness).

Simple example:

```
1  # testing timeit()
2  In [1]: import timeit
3          ...: print(timeit.timeit('output = 30*125'))
4  0.032772750000003015
5  # Combining multiple lines with ;
6  In [2]: print("Summation takes ",timeit.timeit(stmt='x=123.3;y=0.245;sum=x+y'))
7  Summation takes  0.066319874999999781
```


Timeit: measure execution time for code snippets

The Timeit module provides a simple way to time small bits of Python code (it is highly accurate and runs code statements multiple times for robustness).

Simple example:

```
1  # testing timeit()
2  In [1]: import timeit
3          ...: print(timeit.timeit('output = 30*125'))
4  0.032772750000003015
5  # Combining multiple lines with ;
6  In [2]: print("Summation takes ",timeit.timeit(stmt='x=123.3;y=0.245;sum=x+y'))
7  Summation takes  0.066319874999999781
```

Timeit's CLI

Timeit also has a convenient command line interface:

```
1  alex@mac% python -m timeit 'try: ' ' str.__bool__' 'except AttributeError: ' ' pass'
2  1000000 loops, best of 5: 304 nsec per loop
3
4  alex@mac% python -m timeit 'if hasattr(str, "__bool__"): pass'
5  2000000 loops, best of 5: 197 nsec per loop
```

Another use case: timestamps

```
1  import numpy as np
2  import timeit
3
4  N=2**5
5  print("Running test with N =", str(N))
6
7  start = timeit.default_timer()
8  calltime=0
9
10 def fib(n):
11     global calltime
12     calltime+=1
13     if n < 2:
14         return n
15     return fib(n-1) + fib(n-2)
16
17 fibo_list=[fib(n) for n in range(N)]
18 stop = timeit.default_timer()
19
20 print(fibo_list)
21 print("Method recursion: ", round(stop - start,6), ' seconds')
22 print("fib was called ", calltime, " times!")
```

Output from this script

```
1 Running test with N = 32
2 [0, 1, 1, 2, 3, 5, 8, 13, 21, 34, 55, 89, 144, 233, 377, 610, 987, 1597, 2584, 4181,
3 6765, 10946, 17711, 28657, 46368, 75025, 121393, 196418, 317811, 514229, 832040, 1346269]
4 Method recursion: 1.602827 seconds
5 fib was called 11405740 times!
```

Quiz: How can we improve the speed of Fibonacci?

This recursive code is highly inefficient ...

```
1  def fib(n):  
2      if n < 2:  
3          return n  
4      return fib(n-1) + fib(n-2)
```

Let's illustrate this for fib(6):

Reminder: Simple recursion implementation

```
1  N=2**5
2  print("Running test with N =", str(N))
3  start = timeit.default_timer()
4  calltime=0
5  def fib(n):
6      global calltime
7      calltime+=1
8      if n < 2:
9          return n
10     return fib(n-1) + fib(n-2)
11
12  fibo_list=[fib(n) for n in range(N)]
13  stop = timeit.default_timer()
14  print(fibo_list)
15  print("Method recursion: ", round(stop - start,6), ' seconds')
16  print("fib was called ", calltime, " times!")
```

```
1  Running test with N = 32
2  [0, 1, 1, 2, 3, 5, 8, 13, 21, 34, 55, 89, 144, 233, 377, 610, 987, 1597, 2584, 4181, 6765]
3  Method recursion: 1.602827 seconds
4  fib was called 11405740 times!
```


Simple memoization implementation

```
1  start = timeit.default_timer()
2  calltime=0
3  Fibs={0: 0, 1:1}                # initialize memory as dictionary
4  def fib(n):
5      global calltime
6      calltime+=1
7      if n in Fibs.keys():         # if result in memory, look it up!
8          return Fibs[n]
9      else:
10         Fibs[n]=fib(n-1) + fib(n-2) # if result not in memory... compute...
11         return Fibs[n]
12
13  fibo_list=[fib(n) for n in range(N)]
14  stop = timeit.default_timer()
15
16  print("Method dictionary cache: ", round(stop - start,6), ' seconds')
17  print("fib was called ", calltime, " times!")
```

```
1  Method dictionary cache: 6.5e-05 seconds
2  fib was called 92 times!
```

Using built-in cache functools for memoization

```
1  from functools import lru_cache
2  start = timeit.default_timer()
3  calltime=0
4  @lru_cache(maxsize=None)
5  def fib(n):
6      global calltime
7      calltime+=1
8      if n < 2:
9          return n
10     return fib(n-1) + fib(n-2)
11
12  fibo_list = [fib(n) for n in range(N)]
13  stop = timeit.default_timer()
14  print("Method cached: ", round(stop - start,6), ' seconds')
15  print("fib was called ", calltime, " times!")
```

```
1  Method cached:  3.4e-05  seconds
2  fib was called  32  times!
```

Here, `@lru_cache` is a decorator to wrap a function with a memoizing callable that saves up to the `maxsize` most recent calls; here LRU stands for *least recently used*.

Questions?

Clustering

Clustering is an important computational task to group data into different groups (called clusters).

Clustering has many applications. You'll encounter it in many different classes in the future.

- genome analysis
- behavioral analysis
- social network analysis
- ...

Scikit-learn has many clustering algorithms implemented. Check them out, if you are interested!

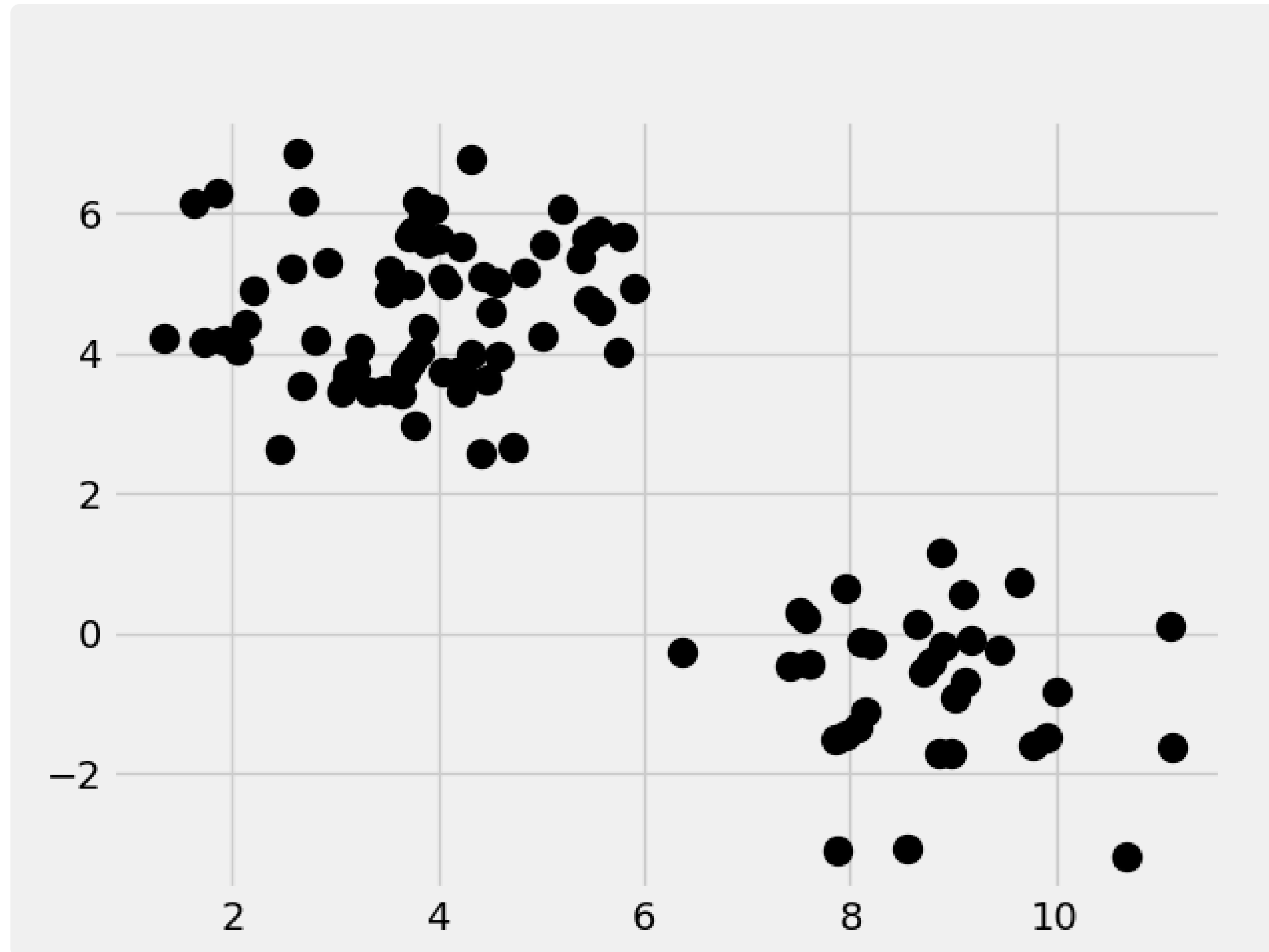
Problem Setting

Imagine you have some data with different features (see the next slide for an example with 2 feature dimensions). The ground truth labels are shown on the slide after that. However, how the data was generated is unknown to us and for clustering, the goal is to assign each data point to one cluster so that the assignment represents the data efficiently.

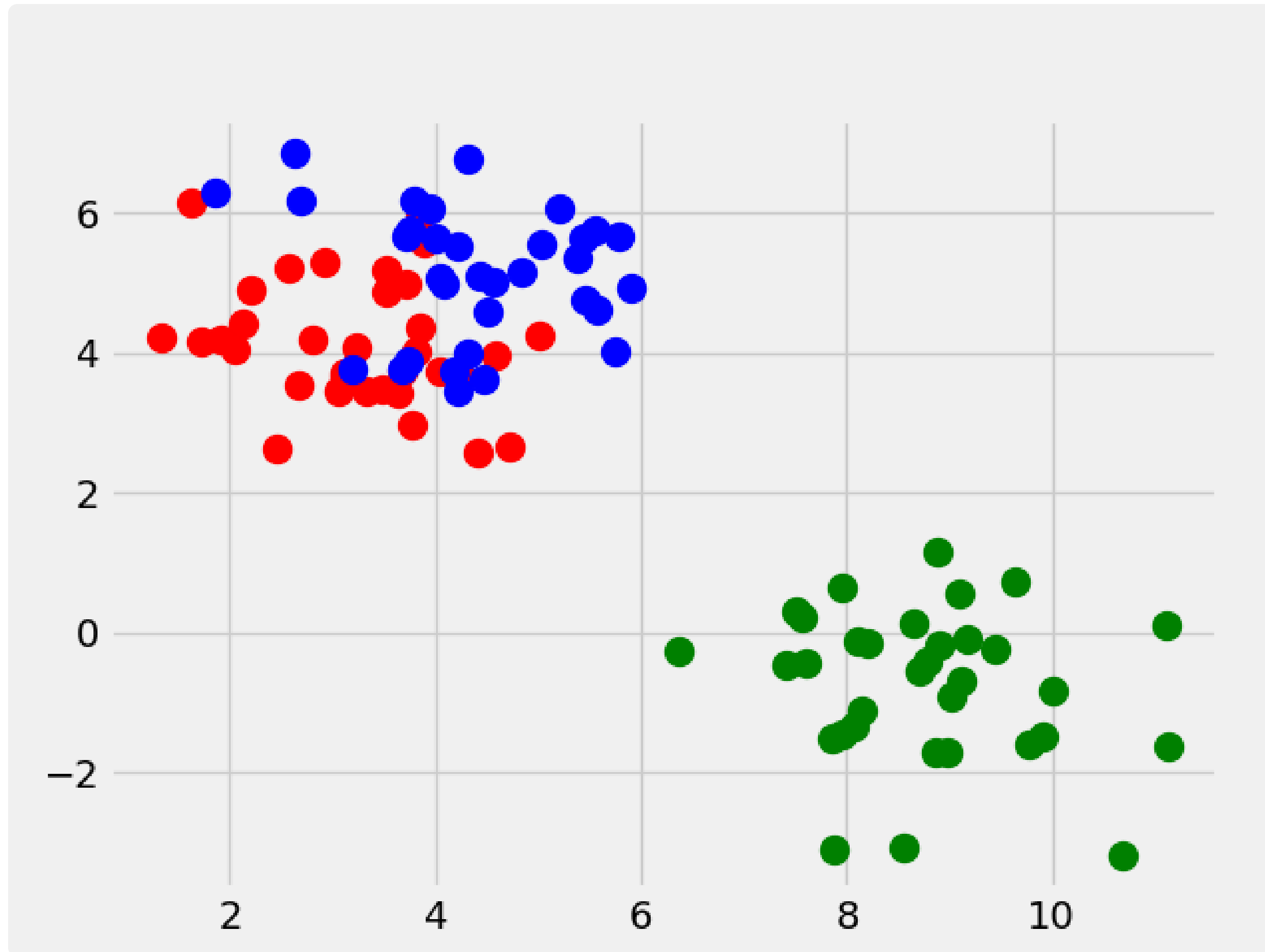
Today we will learn about one classic clustering algorithms, Kmeans. K-means is an iterative algorithm that partitions the data into groups based on the distance between the centroids. The centroid, or cluster center, is the mean of all the points within the cluster.

We pick this algorithm, as it is 1) fundamental algorithm that you should know (you will learn much more in later classes at EPFL) and 2) can nicely be implemented in different ways to show profiling in action.

Example dataset (two features on x and y axis)



Example dataset (ground truth labels)



K-means

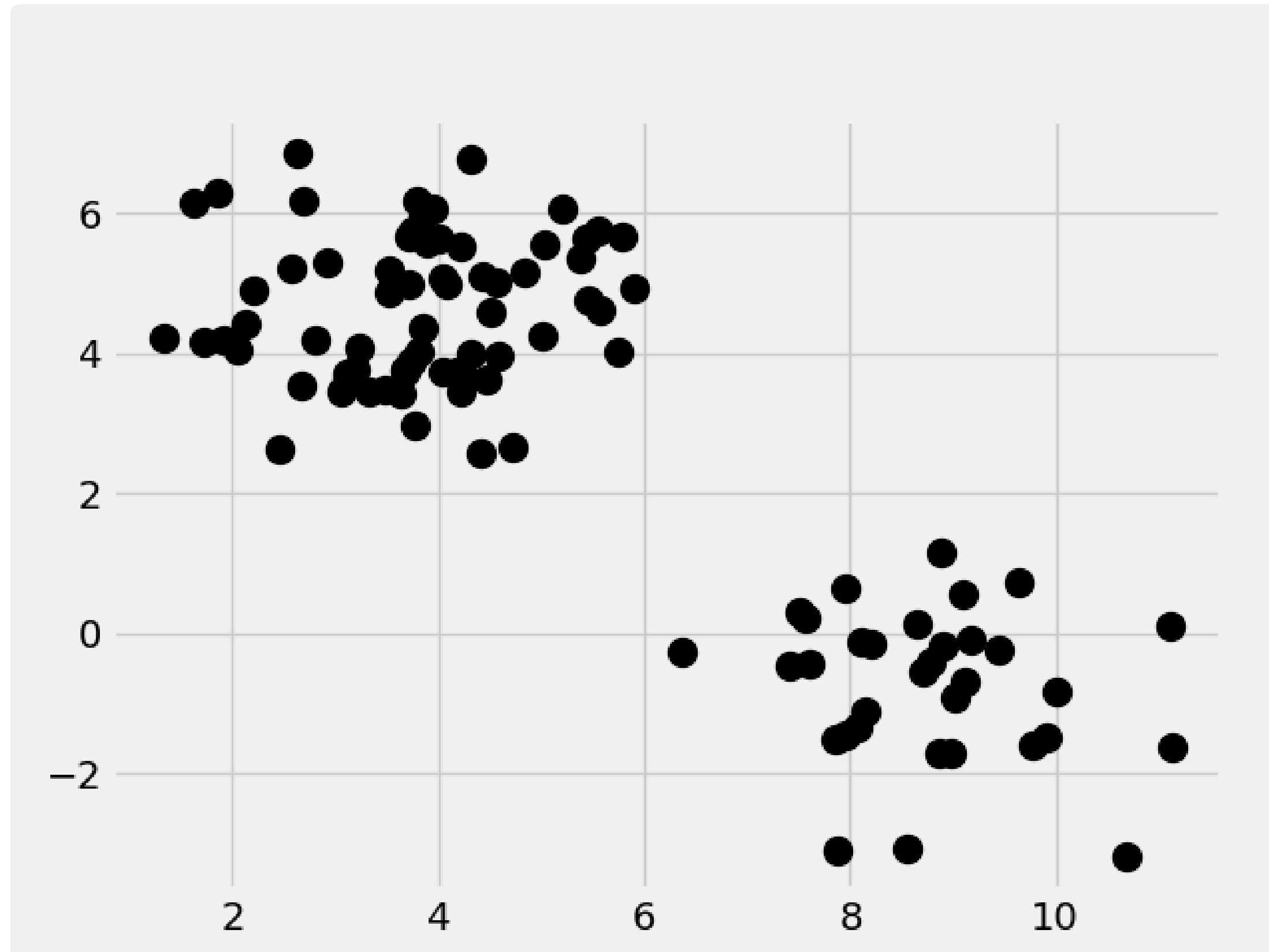
Here is a pseudo algorithm:

First: (randomly) pick centers

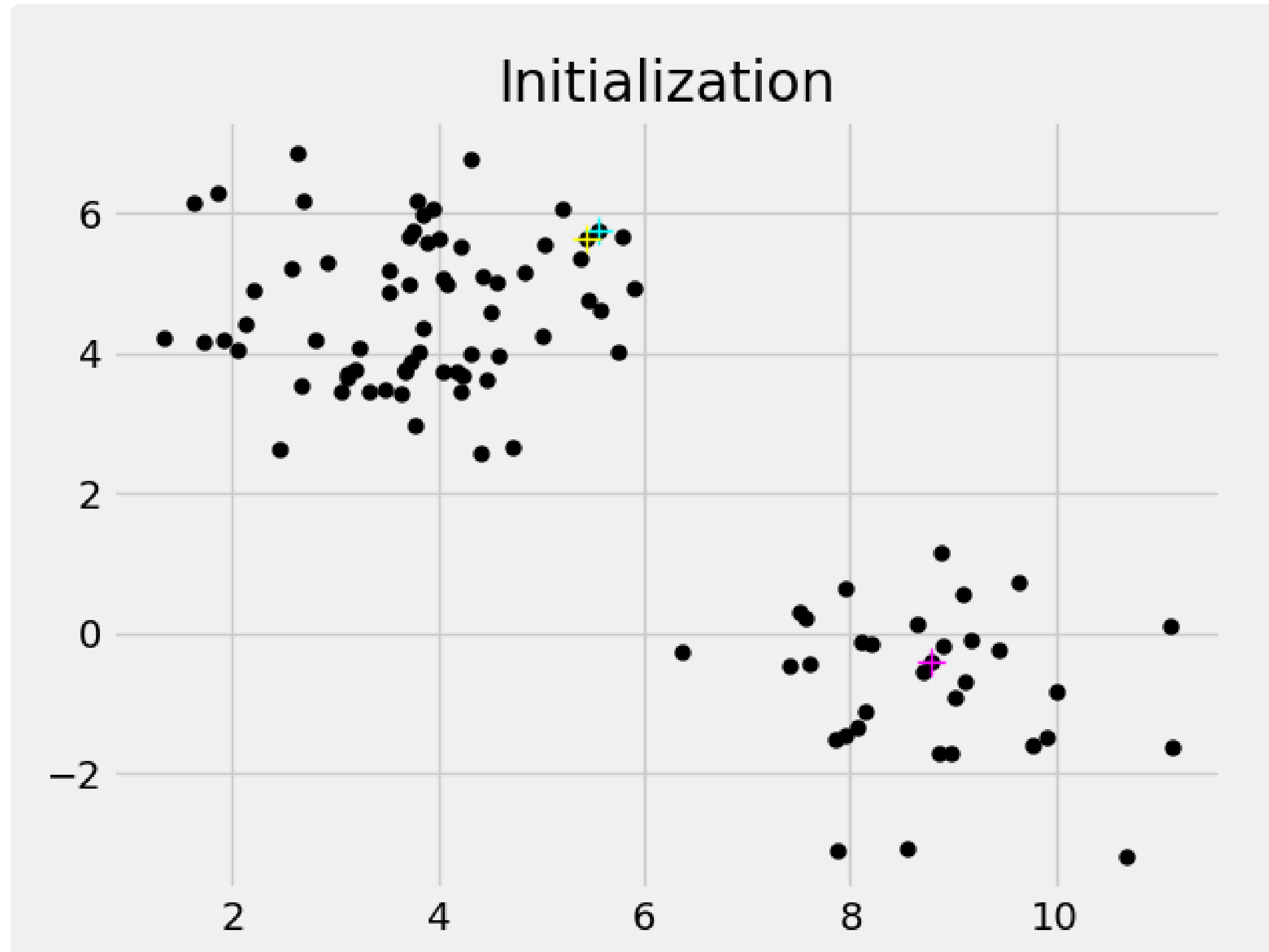
Then iterate:

1. assign data points to nearest center (this defines the clusters)
2. update the center to the mean of the points in each cluster
3. check convergence (stop if converged, otherwise goto 1)

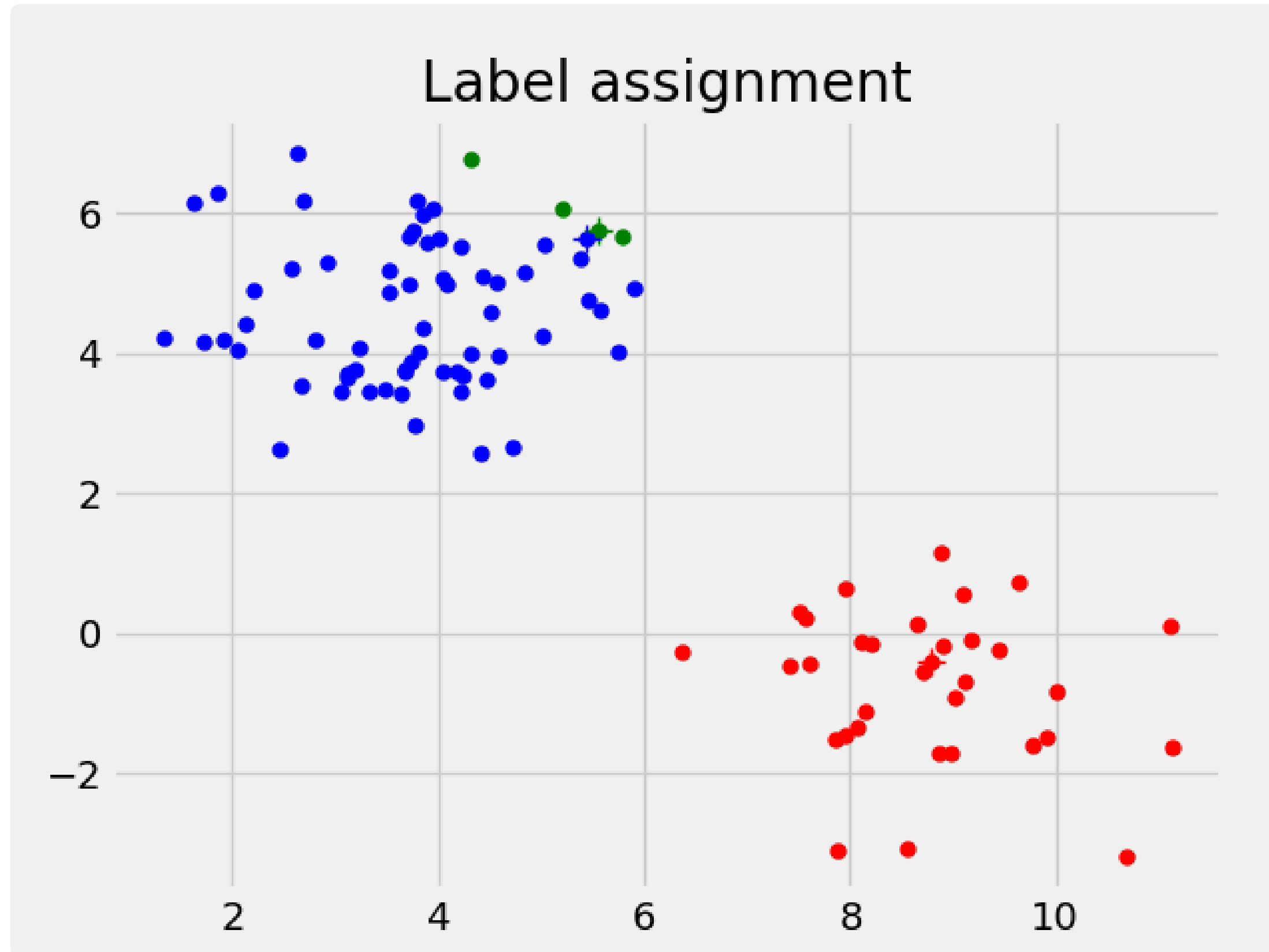
Example dataset



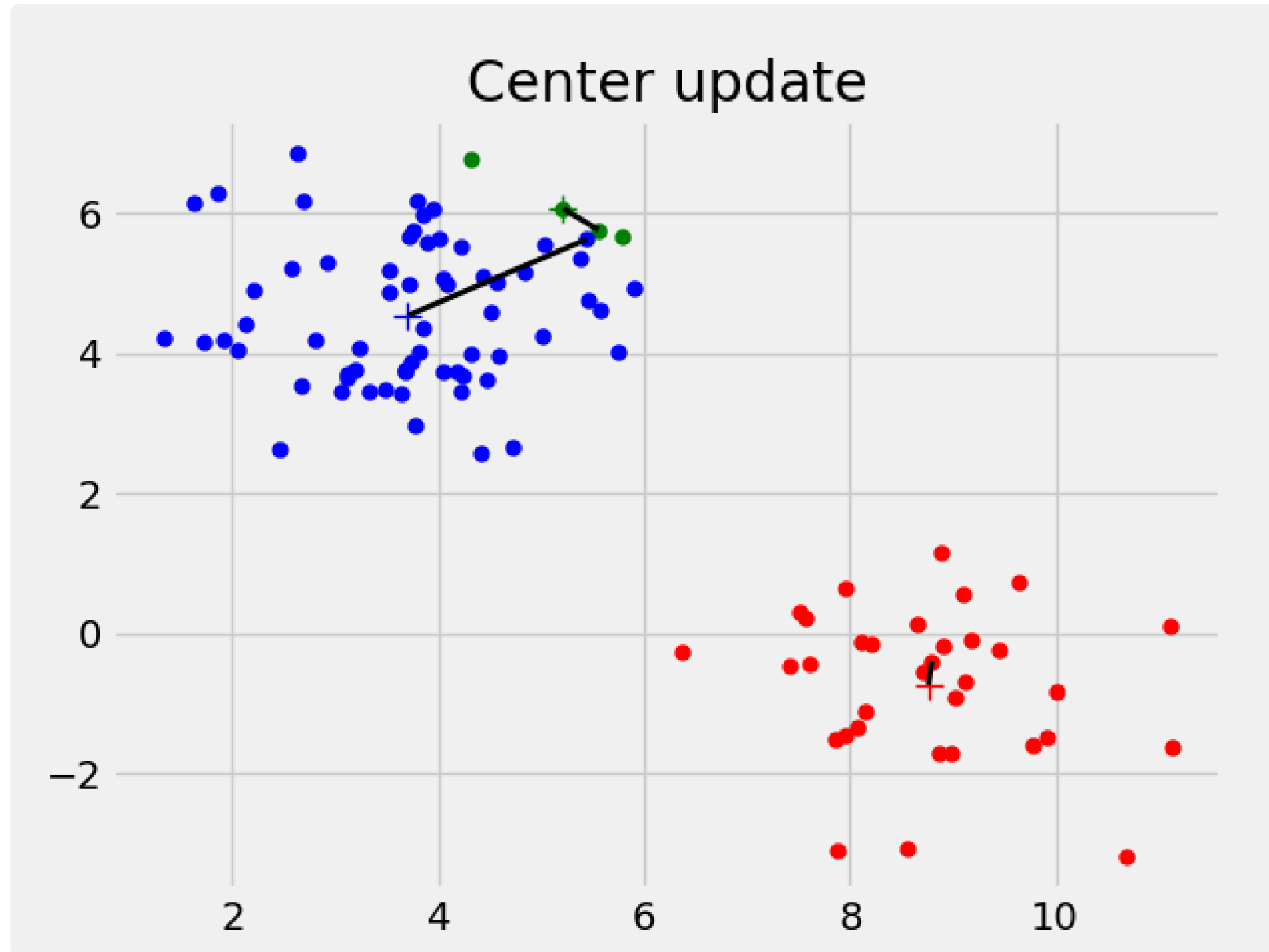
Initialize k-means



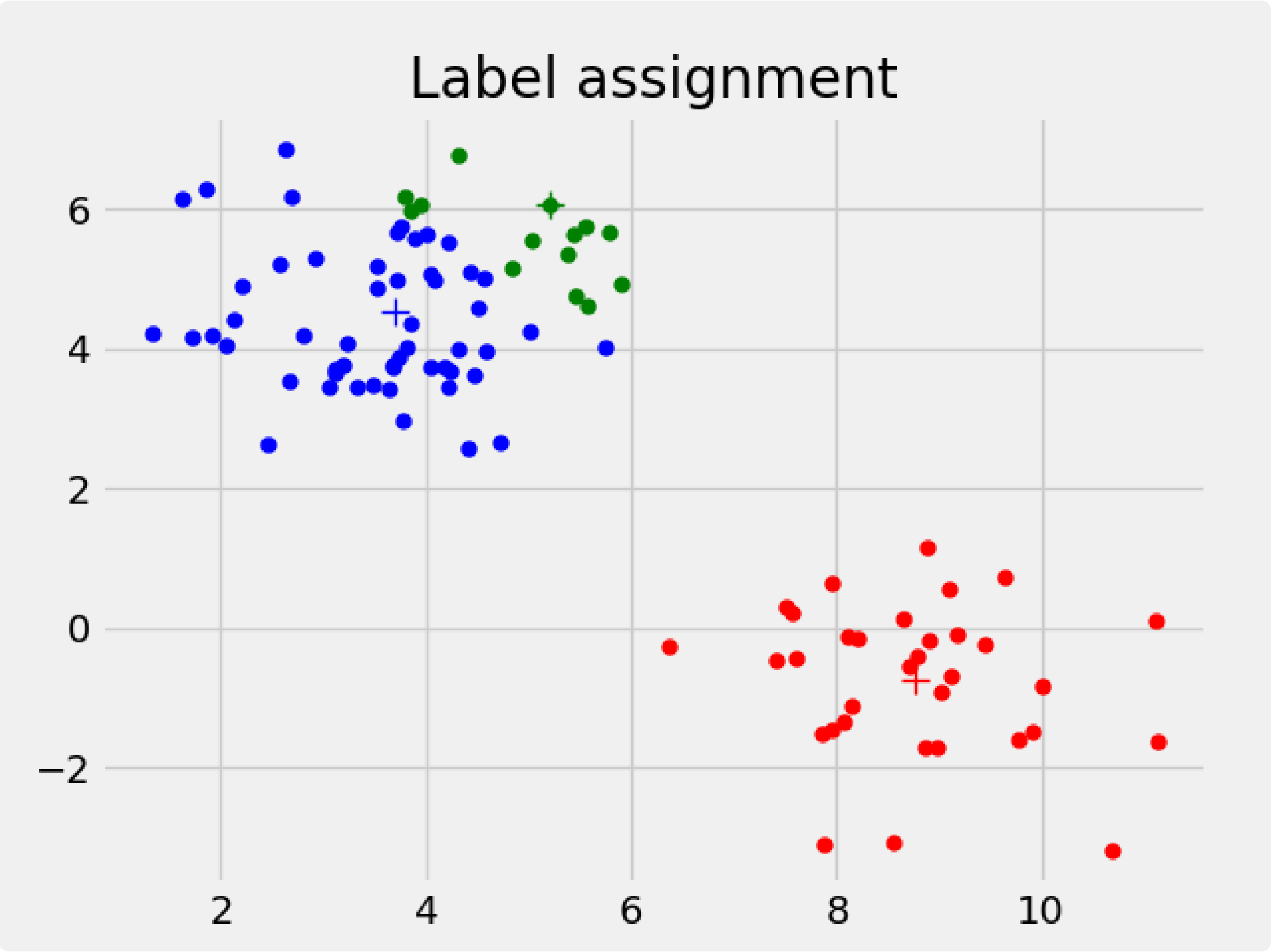
Assign label (iteration: 0)



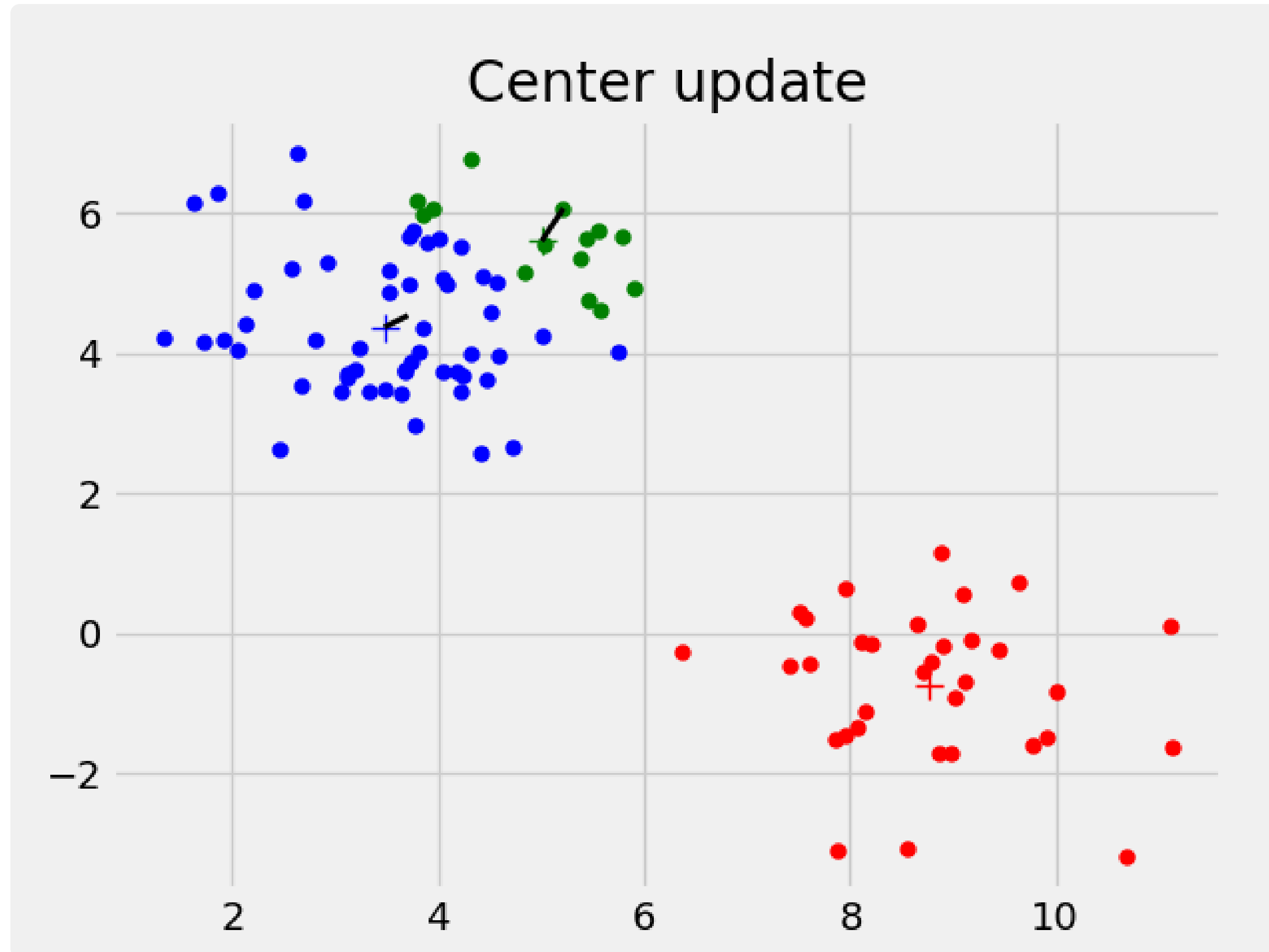
Update centers (iteration: 0)



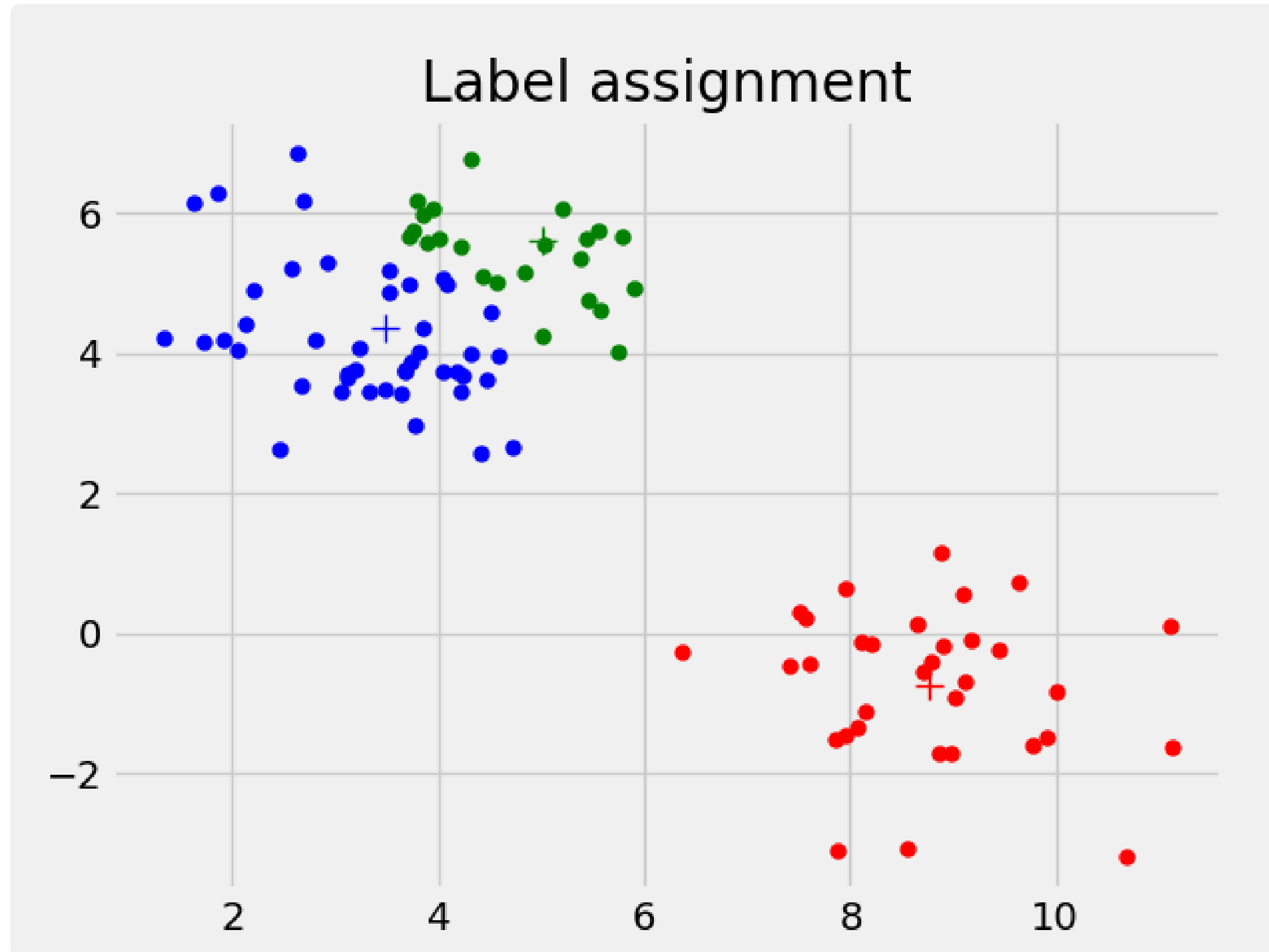
Assign label (iteration: 1)



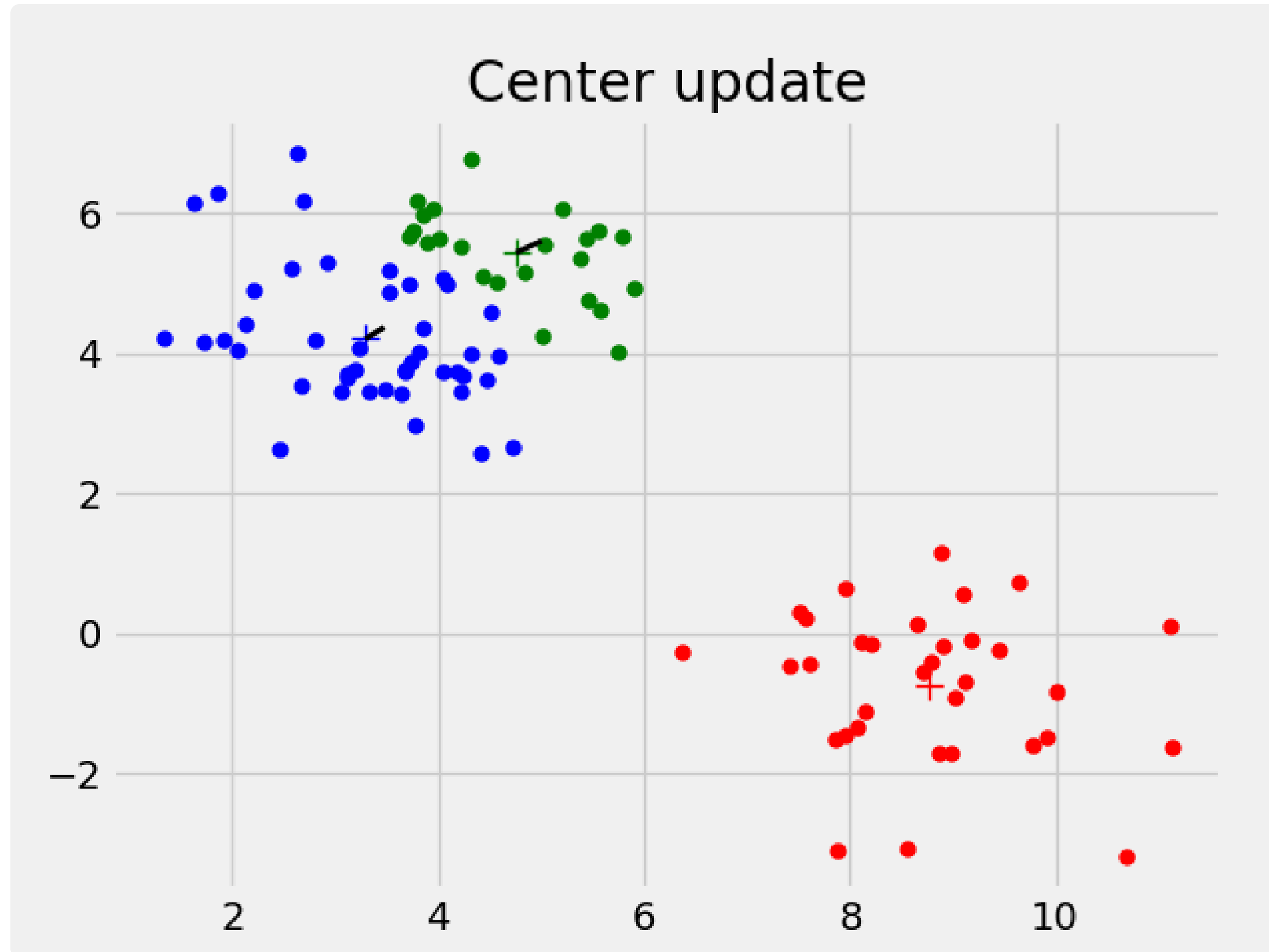
Update centers (iteration: 1)



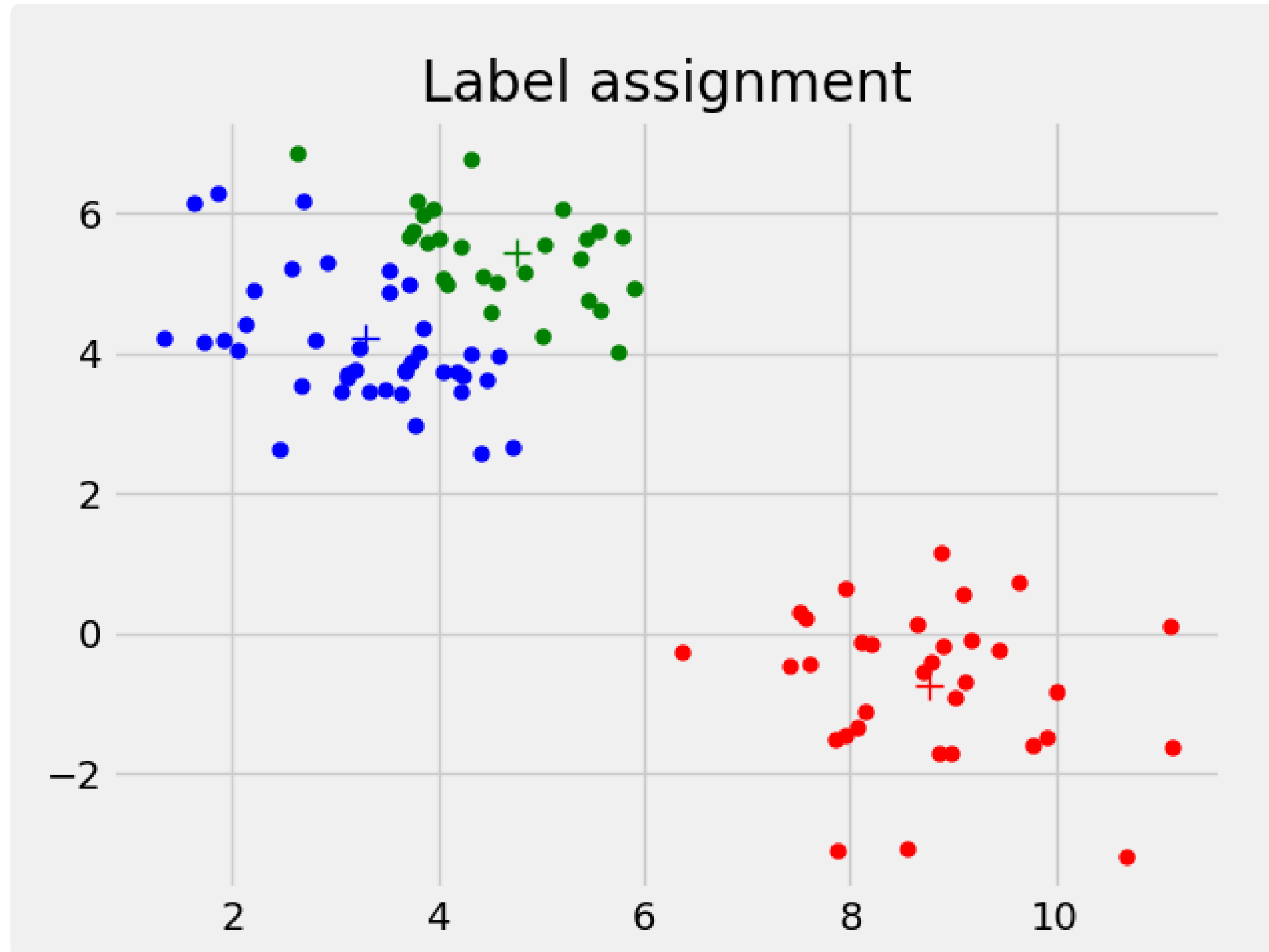
Assign label (iteration: 2)



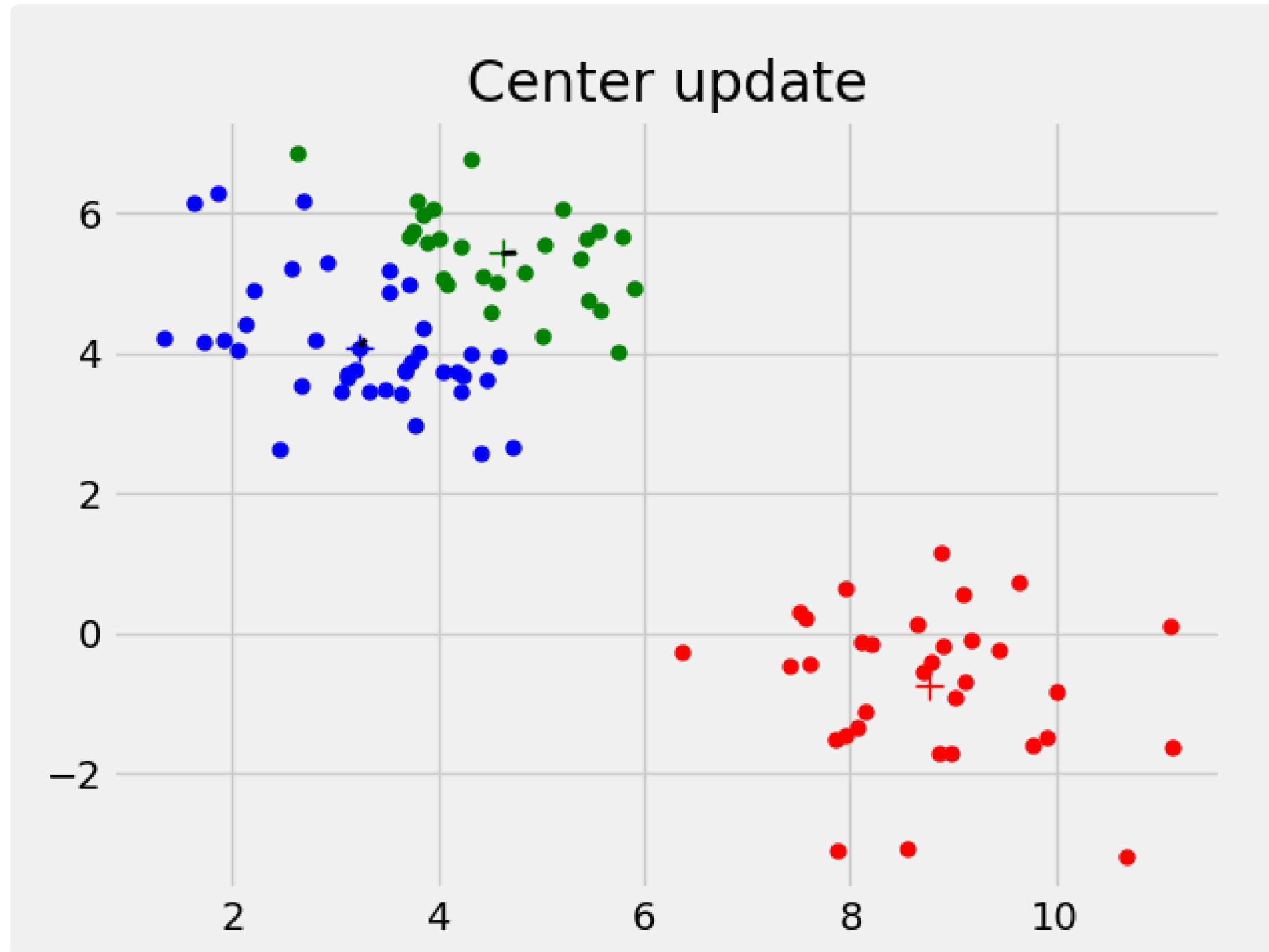
Update centers (iteration: 2)



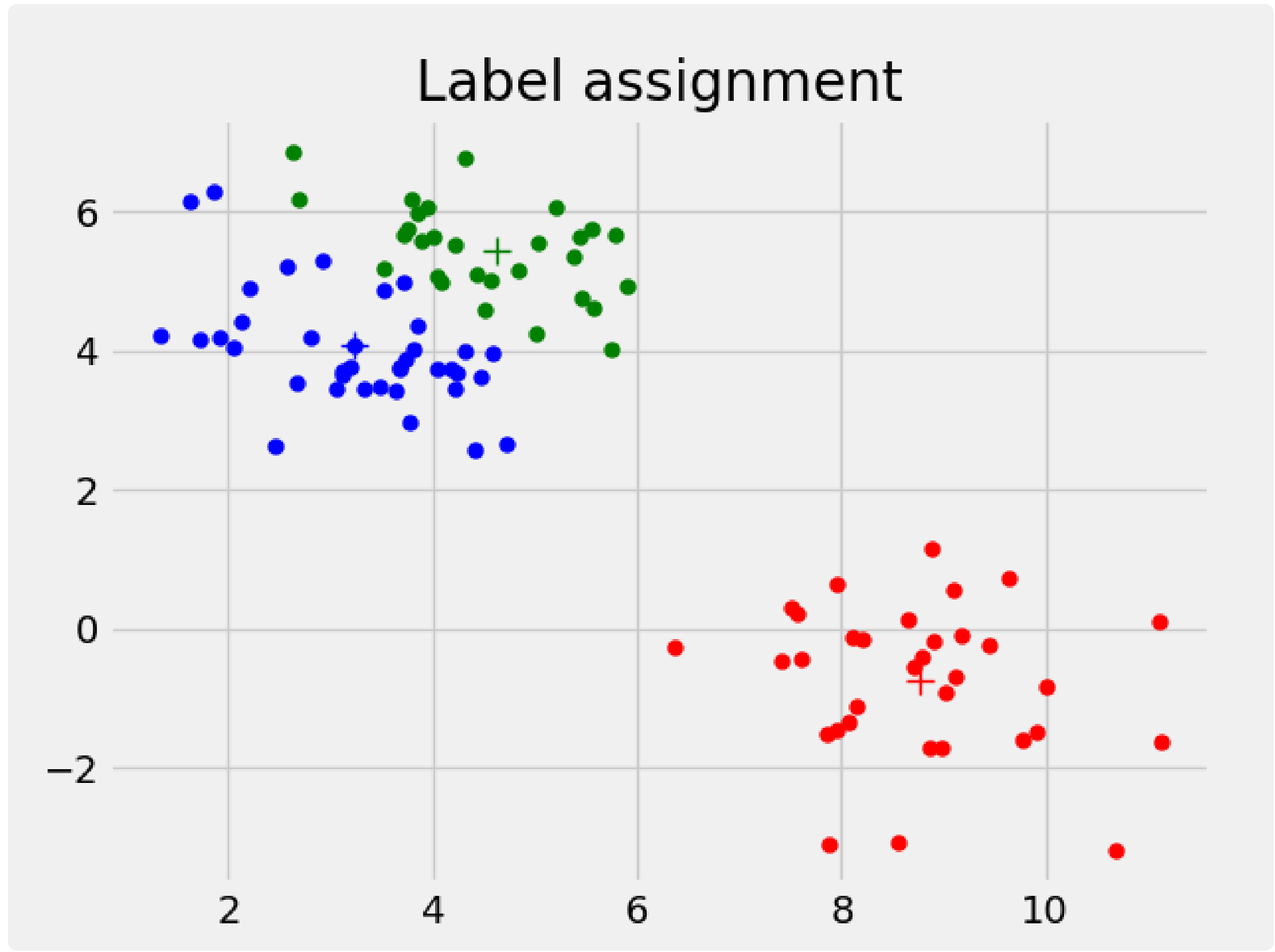
Assign label (iteration: 3)



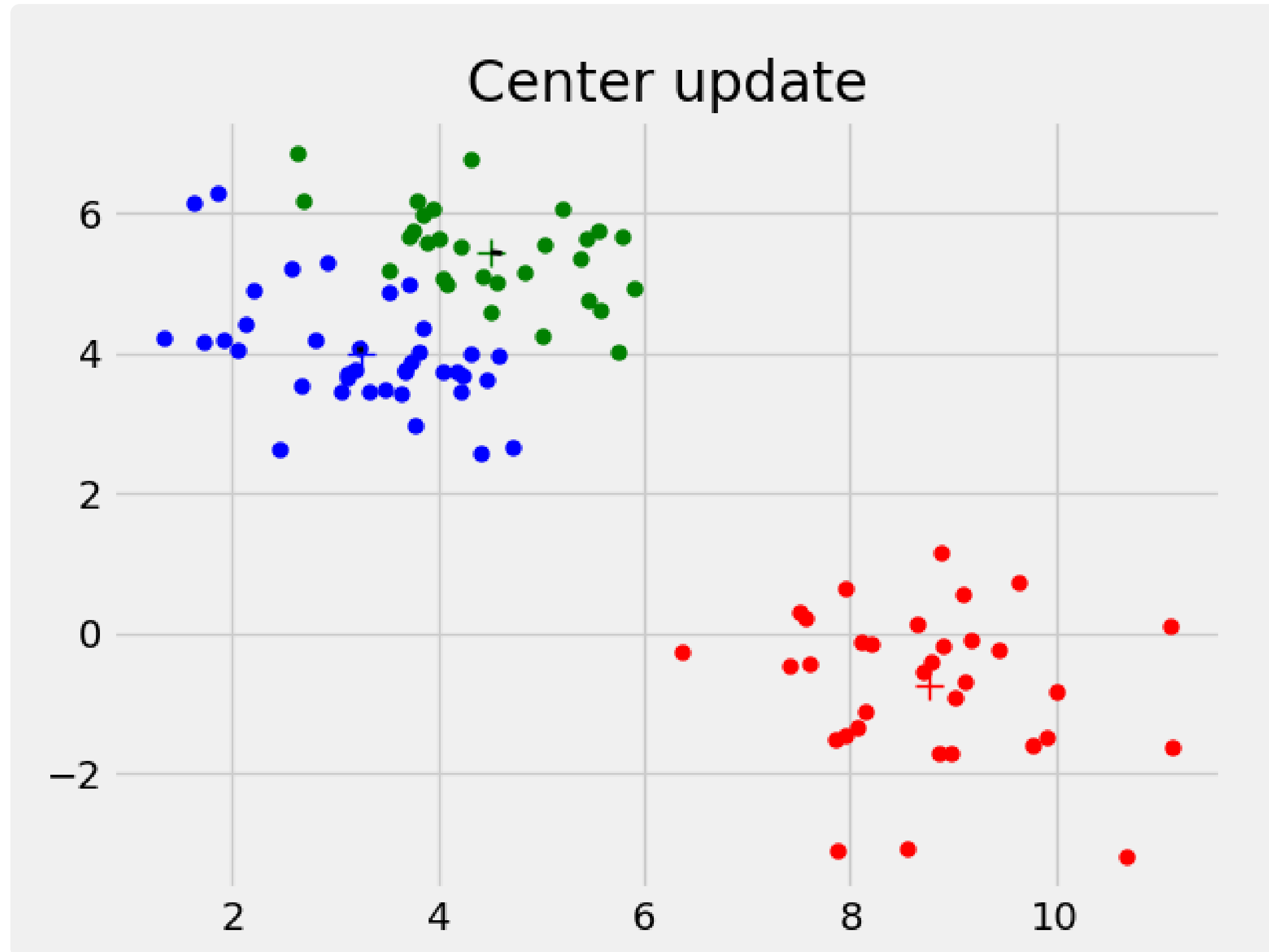
Update centers (iteration: 3)



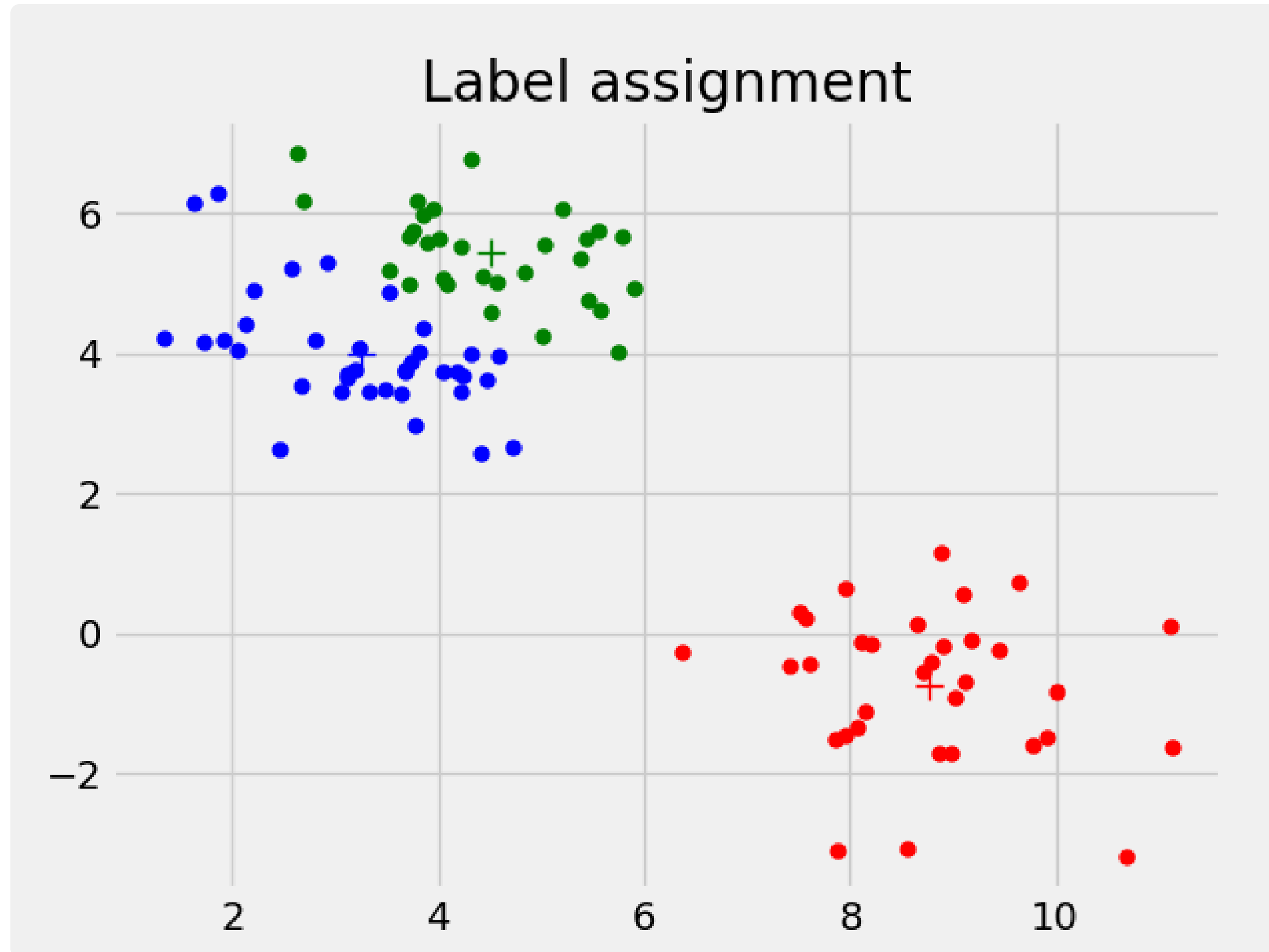
Assign label (iteration: 4)



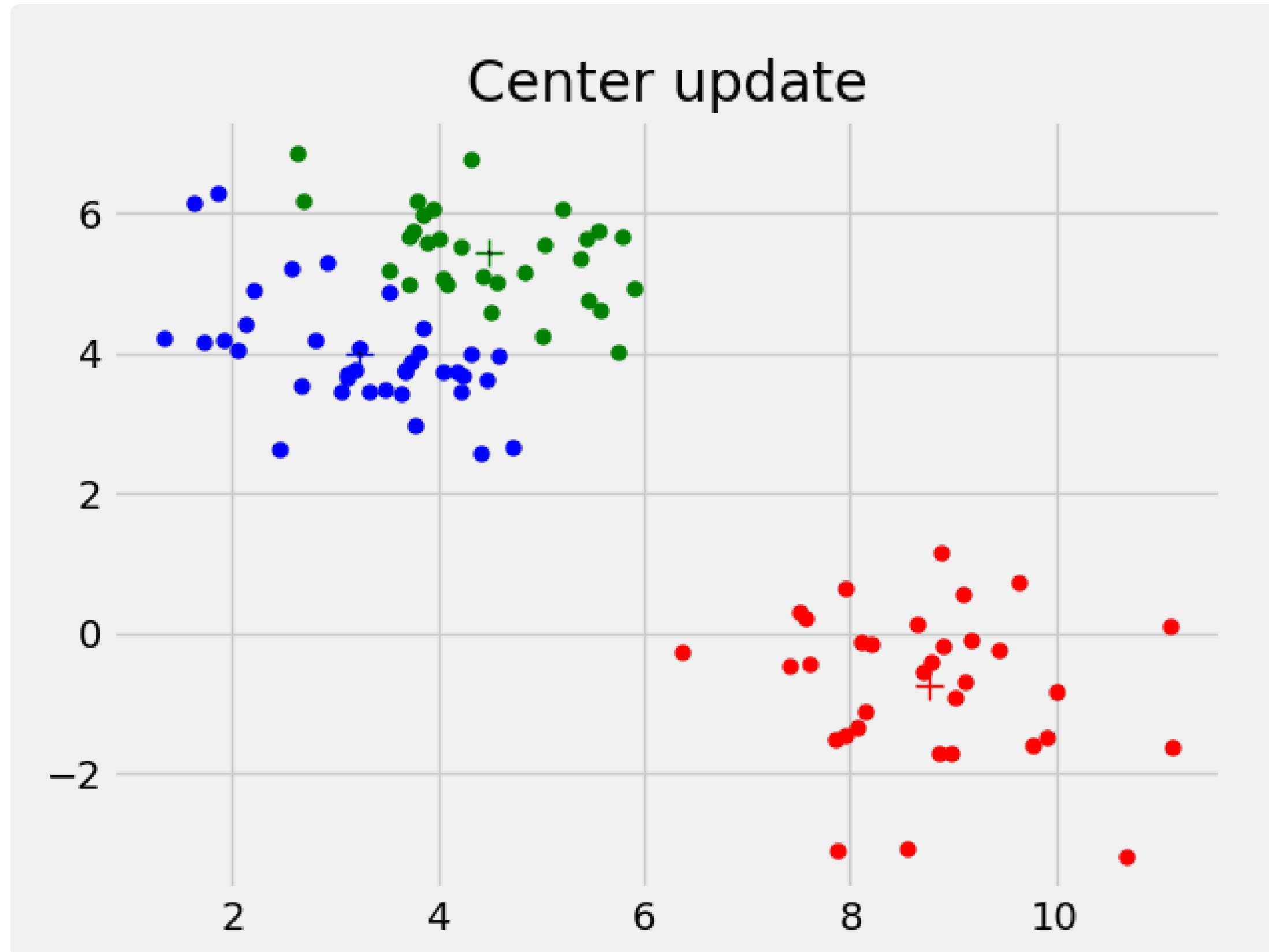
Update centers (iteration: 4)



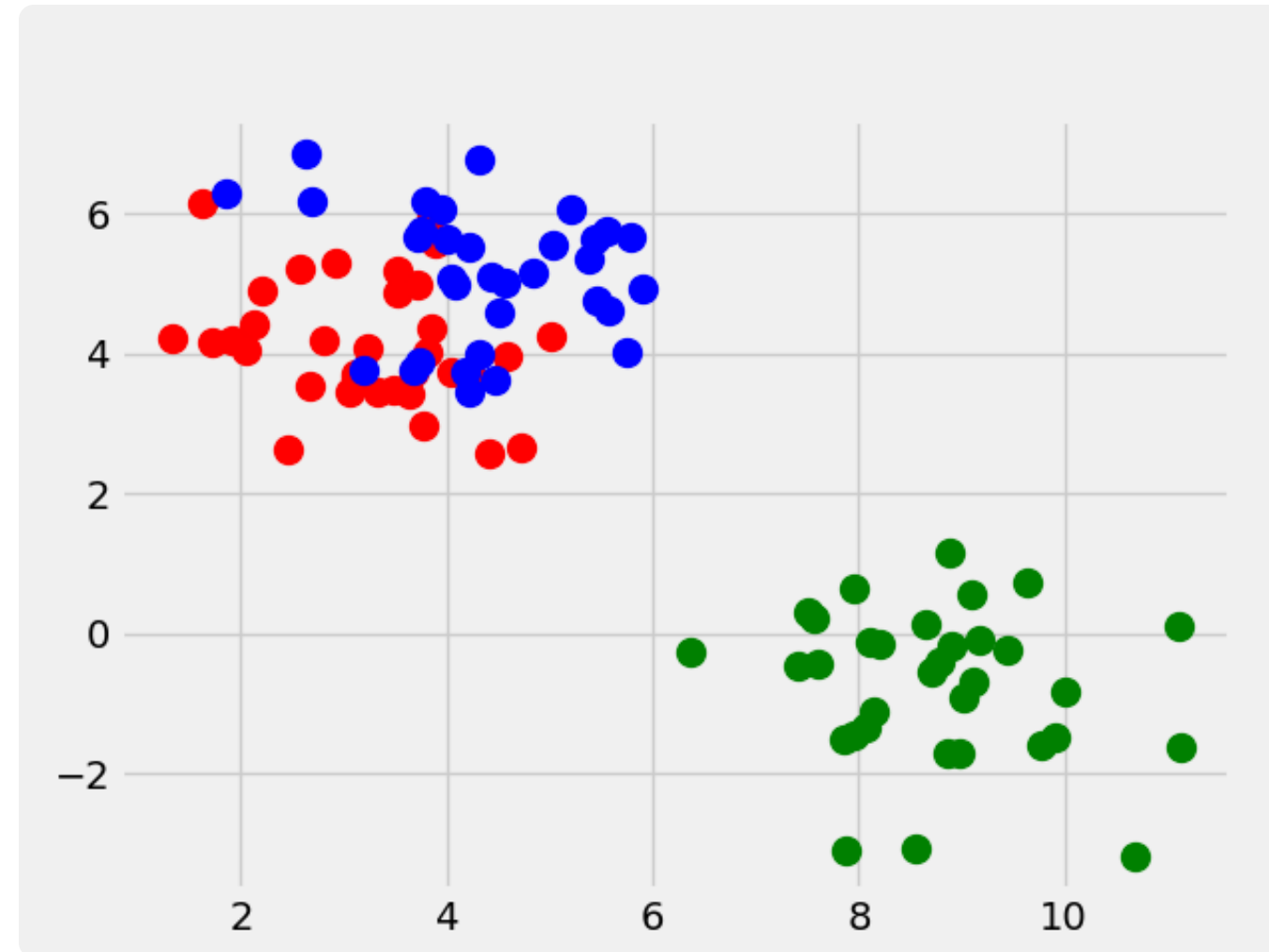
Assign label (iteration: 5)



Update centers (iteration: 5); it converged



Example dataset (ground truth labels)



Side note: Of course, the specific label (here color) cannot be exactly recovered. K-means recovers three reasonable clusters, but of course fails to accurately assign the datapoints where the blue and red clusters mix. If you are interested in learning more, check out alternative clustering algorithms..

We will now return to our goal, profiling this algorithm!

K-means clustering algorithm

First: (randomly) pick centers

Then iterate:

1. assign data points to nearest center (this defines the clusters)
2. update the center to the mean of the points in each cluster
3. check convergence (stop if converged, otherwise goto 1)

Note: Code below adapted from Jake VanderPlas.

Background: zip / zip*

```
1 In [1]: list(zip((0,1,2),(5,3,6)))      # combing two lists into lists of pairs
2 Out[1]: [(0, 5), (1, 3), (2, 6)]
3 In [2]: list(zip(*[(0,1,2),(3,4,5)])) # Unpacking triples into list of 3 tuples
4 Out[2]: [(0, 3), (1, 4), (2, 5)]
5 In [3]: zip?
6 Init signature: zip(self, /, *args, **kwargs)
7 Docstring:
8 zip(*iterables) --> A zip object yielding tuples until an input is exhausted.
9
10     >>> list(zip('abcdefg', range(3), range(4)))
11     [('a', 0, 0), ('b', 1, 1), ('c', 2, 2)]
```

Pure Python implementation: assigning labels

```
1  import math
2  def distance(v, w):
3      """ Computes the Euclidean distance for vectors v and w """
4      return math.sqrt(sum((vi - wi) ** 2 for vi, wi in zip(v, w)))
5
6
7  def assign_labels(data, centers):
8      """ Assign data to closest label (corresponding to centers). """
9      labels = []
10     for item in data:
11         distances = [distance(item, center) for center in centers]
12         labels.append(distances.index(min(distances)))
13     return labels
```

```
1  def update_centers(data, labels):
2      """ Update centers of mass according to labels """
3      n_centers = len(set(labels))
4      n_dims = len(data[0])
5      centers = [[0 for i in range(n_dims)] for j in range(n_centers)]
6      counts = [0 for j in range(n_centers)]
7
8      for label, item in zip(labels, data):
9          counts[label] += 1      # keeping count of # elements and summing per label:
10         centers[label] = [ci + ii for ci, ii in zip(centers[label], item)]
11
12     return [[x / count for x in center] for center, count in zip(centers, counts)]
13
14 def kmeans(data, n_clusters):
15     centers = data[-n_clusters:]
16     iteration = 0
17     while True:
18         old_centers = centers
19         labels = assign_labels(data, centers)
20         centers = update_centers(data, labels)
21
22         if centers == old_centers:
23             break
24
25     return labels
```

Setting up cProfile

```
1  ## Experiment:
2  import kmeans.utils, timeit
3
4  # Creating some data:
5  n_clusters = 10
6  n_dims = 10000
7  data, gt_labels = kmeans.utils.sampledata(n_dims, n_clusters, plotting=False)
8
9  from kmeans.kmeans_python import kmeans
10
11 import cProfile
12
13 cProfile.run("kmeans(data,n_clusters)", "summary_kmeans.stats")
14
15 import pstats
16 stats = pstats.Stats("summary_kmeans.stats")
17 stats.sort_stats("calls")
18 stats.print_stats()
```


Output of cProfile

```
1 Thu Nov 14 18:54:52 2024 summary_kmeans.stats

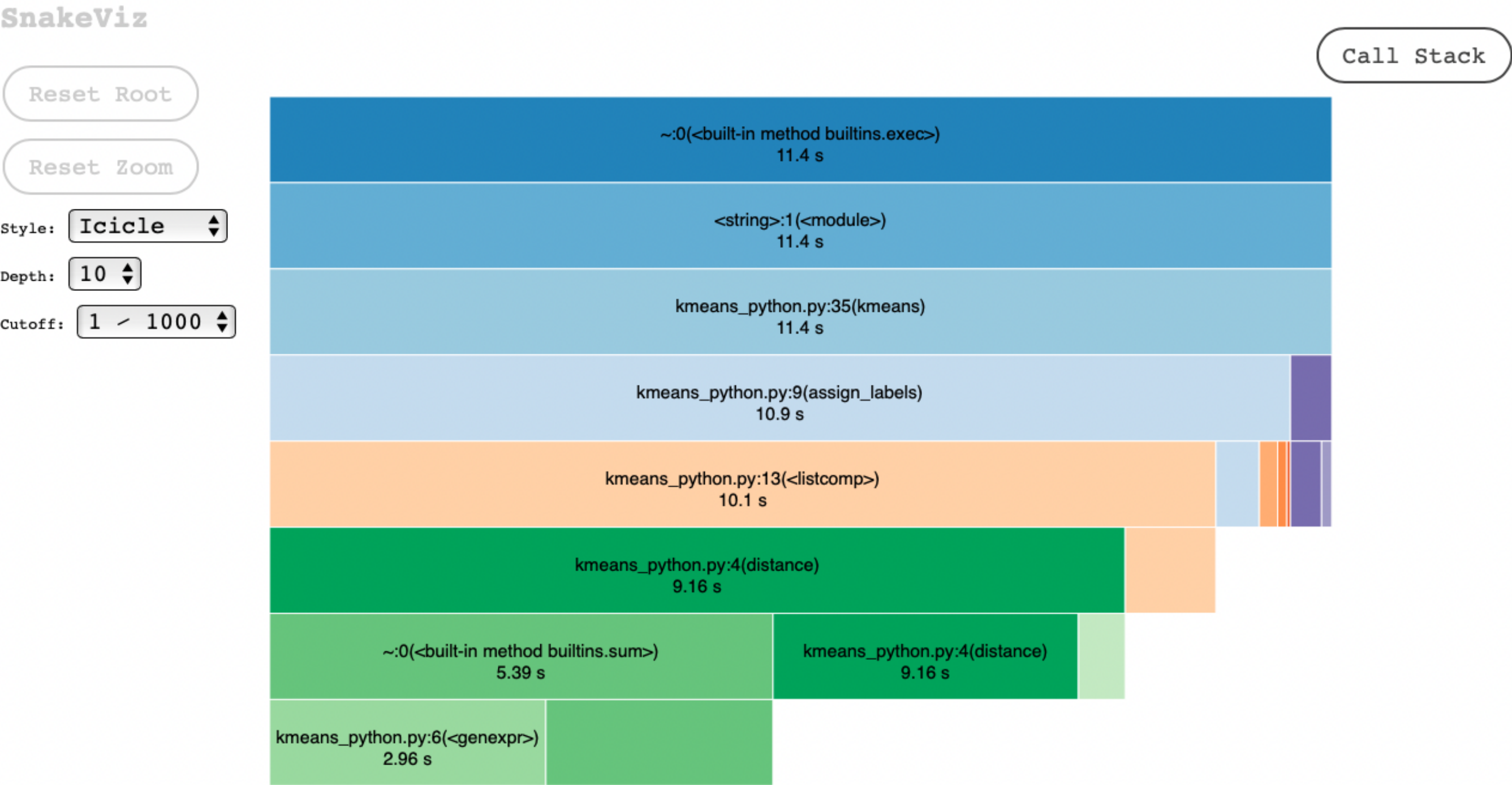
50701252 function calls in 11.380 seconds

Ordered by: call count

ncalls  tottime  percall  cumtime  percall  filename:lineno(function)
23400000  2.959    0.000    2.959    0.000  /Users/alex/Code/Teaching/profiling/kmeans/kmeans.py:100(  
7800000  2.433    0.000    5.392    0.000  {built-in method builtins.sum}  
7800000  0.504    0.000    0.504    0.000  {built-in method math.sqrt}  
7800000  3.269    0.000    9.165    0.000  /Users/alex/Code/Teaching/profiling/kmeans/kmeans.py:100(  
780000  0.041    0.000    0.041    0.000  {method 'append' of 'list' objects}  
780000  0.096    0.000    0.096    0.000  {method 'index' of 'list' objects}  
780000  0.199    0.000    0.199    0.000  {built-in method builtins.min}  
780000  0.971    0.000   10.136    0.000  /Users/alex/Code/Teaching/profiling/kmeans/kmeans.py:100(  
780000  0.112    0.000    0.112    0.000  /Users/alex/Code/Teaching/profiling/kmeans/kmeans.py:100(  
780  0.000    0.000    0.000    0.000  /Users/alex/Code/Teaching/profiling/kmeans/kmeans.py:100(  
156  0.000    0.000    0.000    0.000  {built-in method builtins.len}  
78  0.462    0.006   10.934    0.140  /Users/alex/Code/Teaching/profiling/kmeans/kmeans.py:100(  
78  0.000    0.000    0.000    0.000  /Users/alex/Code/Teaching/profiling/kmeans/kmeans.py:100(  
78  0.000    0.000    0.000    0.000  /Users/alex/Code/Teaching/profiling/kmeans/kmeans.py:100(  
78  0.331    0.004    0.445    0.006  /Users/alex/Code/Teaching/profiling/kmeans/kmeans.py:100(  
1  0.000    0.000   11.380   11.380  {built-in method builtins.exec}  
1  0.000    0.000    0.000    0.000  {method 'disable' of '_lsprof.Profiler' objects}
```

Visualization type 1 (icicle)

```
1 snakeviz summary_kmeans.stats
```



Visualization type 2 (sunburst)

```
1 snakeviz summary_kmeans.stats
```

SnakeViz

Reset Root

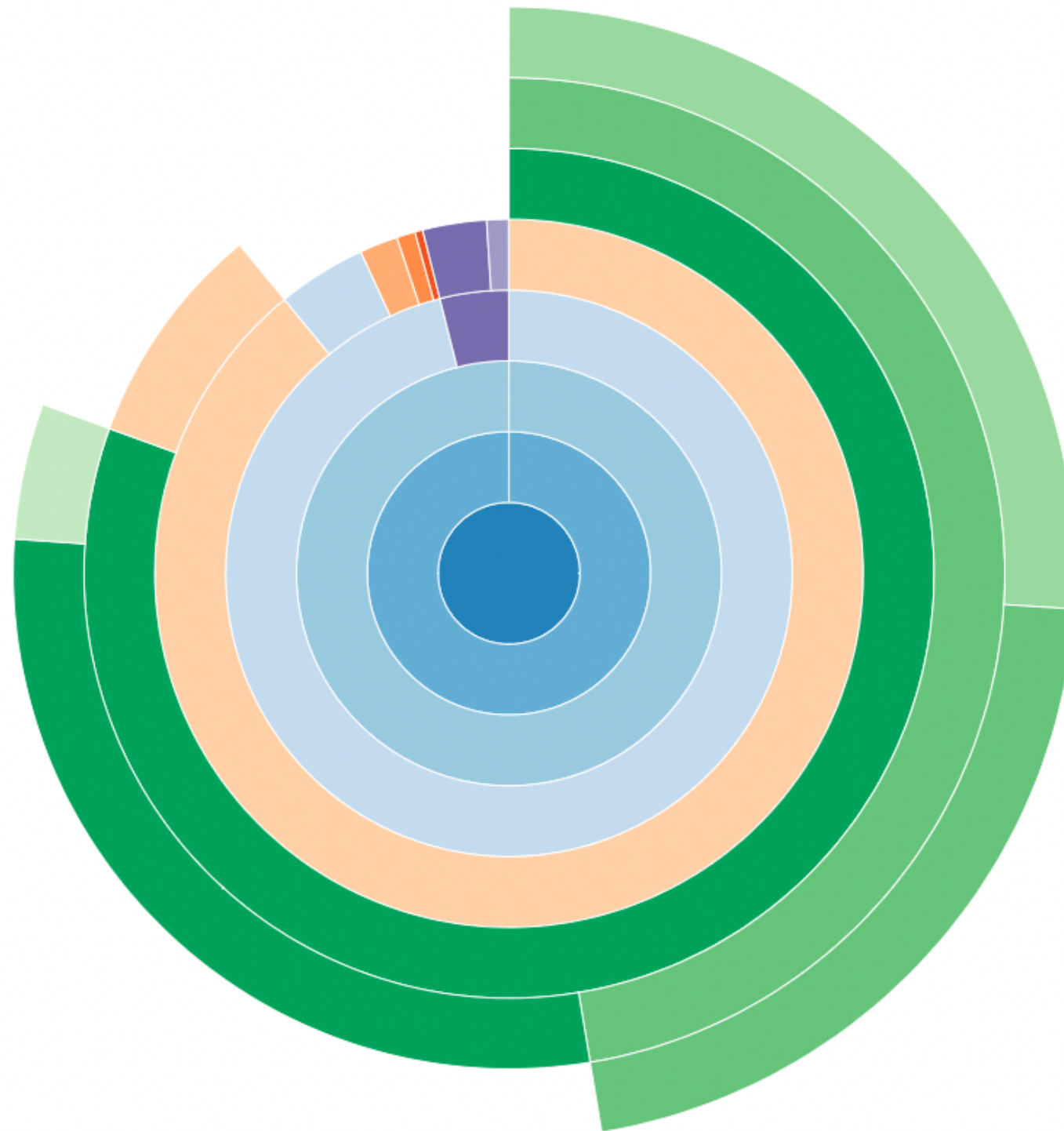
Reset Zoom

Style: **Sunburst** ▾

Depth: **10** ▾

Cutoff: **1 / 1000** ▾

Call Stack



Alternative profiling (line_profiling)

Especially for numerical code, from the cProfile output it is

- not apparent what exactly is computed (without going back to the code)
- calls/times are grouped by function calls

What if much work happens in particular lines, e.g. loading data, diagonalizing a matrix ...

For instance, most calls are from this function:

```
1  ncalls  tottime  percall  cumtime  percall filename:lineno(function)
23400000  2.959    0.000    2.959    0.000 /Users/alex/Code/Teaching/profiling/kmeans/kmeans.py
```

Line-by-line profiling

`line_profiler` is a module for doing line-by-line profiling of functions. kernprof is a convenient script for running either line_profiler.

Here we will use it in ipython with magic commands

Installation:

```
1 pip install line_profiler
```

Demo code (LineProfiling.py)

```
1  import utils
2
3  n_clusters=10      # Initalizing data
4  n_dims =10000
5  data, gt_labels = utils.sampledata(n_dims, n_clusters, plotting= False)
6
7  from kmeans_python import kmeans
8  labels = kmeans(data, n_clusters)    # Running our kmeans
9  utils.score(gt_labels,labels)        # Evaluating performance
```

```

1 In [1]: run LineProfiling.py
2 0.9191471147114711
3 In [2]: %load_ext line_profiler
4 In [3]: %lprun -f kmeans kmeans(data,n_clusters)
5 Timer unit: 1e-06 s
6 Total time: 8.55269 s
7 File: /Users/alex/Code/Teaching/profiling/kmeans/kmeans_python.py
8 Function: kmeans at line 35
9
10 Line #           Hits           Time  Per Hit    % Time  Line Contents
11 =====
12      35                                def kmeans(data, n_clusters):
13      36                1          30.0      30.0      0.0          centers = data[-n_clusters:]
14      37
15      38                1           2.0       2.0      0.0          iteration = 0
16      39              while True:
17      40                41          31.0       0.8      0.0              old_centers = centers
18      41                41      7422540.0  181037.6     86.8              labels = assign_labels(data, cen
19      42                41      1130007.0   27561.1     13.2              centers = update_centers(data, l
20      43
21      44                41          81.0       2.0      0.0              if centers == old_centers:
22      45                1           1.0       1.0      0.0                  break
23      46
24      47                1           0.0       0.0      0.0          return labels

```

We spend >86% in Line 41!!!

Quiz: How can we make this faster?

```
1  def distance(v, w):
2      return math.sqrt(sum((vi - wi) ** 2 for vi, wi in zip(v, w)))
3
4  def assign_labels(data, centers):
5      labels = []
6      for item in data:
7          distances = [distance(item, center) for center in centers]
8          labels.append(distances.index(min(distances)))
9      return labels
```


Numpy to the rescue!

```
1 def distance(v, w):
2     return math.sqrt(sum((vi - wi) ** 2 for vi, wi in zip(v, w)))
3
4 def assign_labels(data, centers):
5     labels = []
6     for item in data:
7         distances = [distance(item, center) for center in centers]
8         labels.append(distances.index(min(distances)))
9     return labels
```

Numpy code:

```
1 def assign_labels(data, centers):
2     """ assign data to closest label (corresponding to centers)."""
3     differences = (np.array(data)[: , None, :] - np.array(centers)) ** 2
4     distances = np.sqrt(differences.sum(-1))
5     return distances.argmin(1)
```

Quick explanation:

```
1 In [1]: import numpy as np
2
3 In [2]: data = np.arange(6).reshape((3,2))
4
5 In [3]: centers = np.arange(2).reshape((1,2))
6
7 In [4]: data[:,None,:] - centers          # introduces extra dimension in middle
8 Out[4]:
9 array([[[0, 0]],
10        [[2, 2]],
11        [[4, 4]]])
12
13
14
15 In [5]: np.sqrt(np.sum((data[:,None,:] - centers)**2,axis = 2))
16 Out[5]:
17 array([[0.          ],          # is distance of (1,2) and (1,2)
18        [2.82842712],          # is distance of (3,4) and (1,2) = sqrt(8)
19        [5.65685425]])          # is the distance of (5,6) and (1,2)
```

Updated line profiling 8.5s -> 🙈

```
1 In [8]: %lprun -f kmeans kmeans(data,n_clusters)
2 Timer unit: 1e-06 s
3
4 Total time: 1.88813 s
5 File: /Users/alex/Code/Teaching/profiling/kmeans/kmeans_numpyI.py
6 Function: kmeans at line 28
7
8 Line #      Hits          Time Per Hit   % Time  Line Contents
9 =====
10      28                               def kmeans(data, n_clusters):
11      29              1           5.0         5.0         0.0      centers = data[-n_clusters:]
12      30
13      31                          while True:
14      32              53          24.0           0.5         0.0          old_centers = centers
15      33              53      334932.0       6319.5       17.7          labels = assign_labels(data, cen
16      34              53      1553147.0      29304.7       82.3          centers = update_centers(data, l
17      35              53          27.0           0.5         0.0          if centers == old_centers:
18      36              1           0.0           0.0         0.0              break
19      37
20      38              1           0.0           0.0         0.0      return labels
```

Now we spend >80% for updating centers!

Quiz: How can we improve this code?

```
1  def update_centers(data, labels):
2      """ Update centers of mass according to labels """
3      n_centers = len(set(labels))
4      n_dims = len(data[0])
5
6      centers = [[0 for i in range(n_dims)] for j in range(n_centers)]
7      counts = [0 for j in range(n_centers)]
8
9      for label, item in zip(labels, data):
10         counts[label] += 1
11         centers[label] = [
12             ci + ii for ci, ii in zip(centers[label], item)
13         ] # summing, per label
14
15     return [[x / count for x in center] for center, count in zip(centers, counts)]
```

Numpy to the rescue, again!

```
1  def update_centers(data, labels):
2      """ Update centers of mass according to labels """
3      n_centers = len((set(labels)))
4      n_dims = len(data[0])
5      centers = [[0 for i in range(n_dims)] for j in range(n_centers)]
6      counts = [0 for j in range(n_centers)]
7      for label, item in zip(labels, data):
8          counts[label] += 1
9          centers[label] = [
10              ci + ii for ci, ii in zip(centers[label], item)
11          ] # summing, per label
12      return [[x / count for x in center] for center, count in zip(centers, counts)]
```

Numpy code:

```
1  def update_centers(data, labels):
2      """ Update centers of mass according to labels """
3      n_centers = len((set(labels)))
4      return np.array([data[labels == i].mean(0) for i in range(n_centers)])
```

Updated line-profiling results... 1.8s --> 🙈

```
1 In [11]: %lprun -f kmeans kmeans(data,n_clusters)
2 Timer unit: 1e-06 s
3
4 Total time: 0.196476 s
5 File: /Users/alex/Code/Teaching/profiling/kmeans/kmeans_numpyII.py
6 Function: kmeans at line 17
7
8 Line #      Hits          Time Per Hit     % Time  Line Contents
9 =====
10      17                                def kmeans(data, n_clusters):
11      18                1      10147.0   10147.0         5.2      data = np.array(data)
12      19                1         10.0        10.0         0.0      centers = data[-n_clusters:]
13      20
14      21                                while True:
15      22                41         17.0         0.4         0.0          old_centers = centers
16      23                41     122246.0    2981.6        62.2          labels = assign_labels(data, cen
17      24                41     63848.0    1557.3        32.5          centers = update_centers(data, l
18      25
19      26                41         207.0         5.0         0.1          if (centers == old_centers).all(
20      27                1          1.0          1.0         0.0              break
21      28
22      29                1          0.0          0.0         0.0      return labels
```

Quiz: What now?

Tip: use specialized functions!

```
1  from scipy.spatial.distance import cdist
2
3  def assign_labels(data, centers):
4      """ assign data to closest label (corresponding to centers)."""
5      distances = cdist(data, centers)
6      return distances.argmin(1)
```

for comparison:

```
1  def assign_labels(data, centers):
2      """ assign data to closest label (corresponding to centers)."""
3      differences = (np.array(data)[: , None, :] - np.array(centers)) ** 2
4      distances = np.sqrt(differences.sum(-1))
5      return distances.argmin(1)
```


Updated line-profiling results... 0.19s --> 🙈

```
1 In [13]: %lprun -f kmeans kmeans(data,n_clusters)
Timer unit: 1e-06 s
```

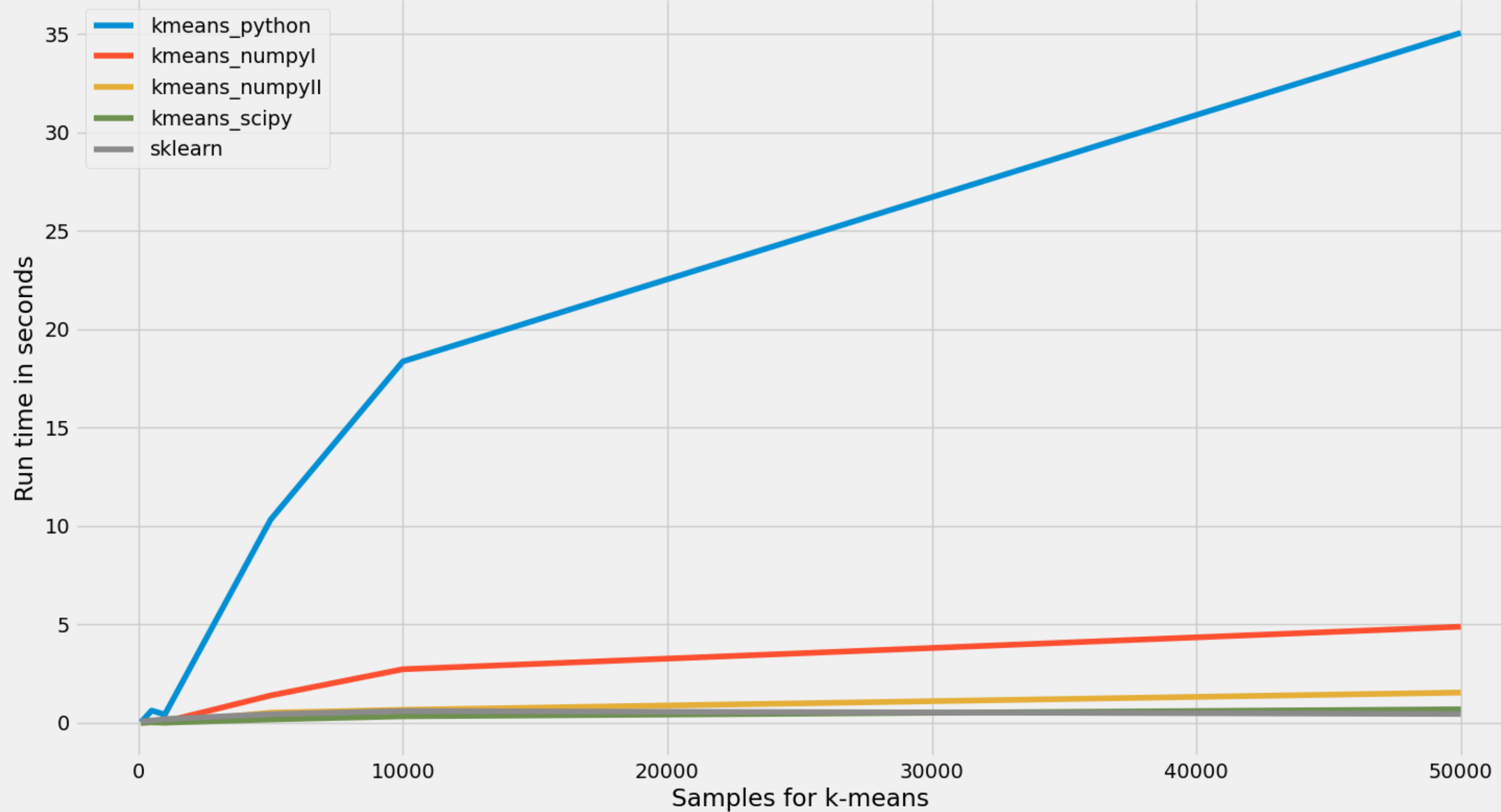
Total time: 0.092516 s

File: /Users/alex/Code/Teaching/profiling/kmeans/kmeans_scipyI.py

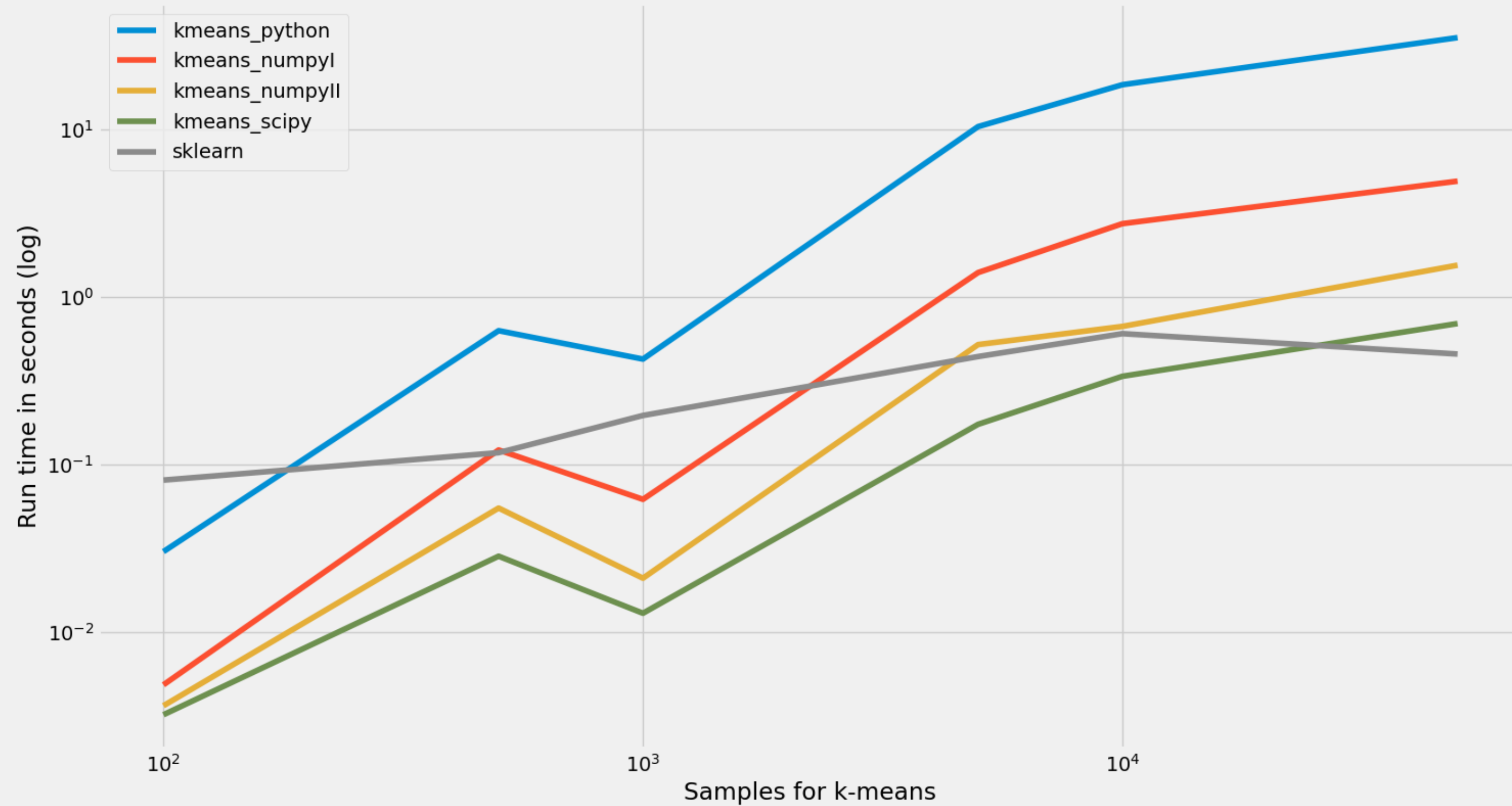
Function: kmeans at line 17

Line #	Hits	Time	Per Hit	% Time	Line Contents
17					def kmeans(data, n_clusters):
18	1	8087.0	8087.0	8.7	data = np.array(data)
19	1	11.0	11.0	0.0	centers = data[-n_clusters:]
20					
21					while True:
22	36	18.0	0.5	0.0	old_centers = centers
23	36	23362.0	648.9	25.3	labels = assign_labels(data, centers)
24	36	60846.0	1690.2	65.8	centers = update_centers(data, labels)
25					
26	36	191.0	5.3	0.2	if (centers == old_centers).all():
27	1	0.0	0.0	0.0	break
28					
29	1	1.0	1.0	0.0	return labels

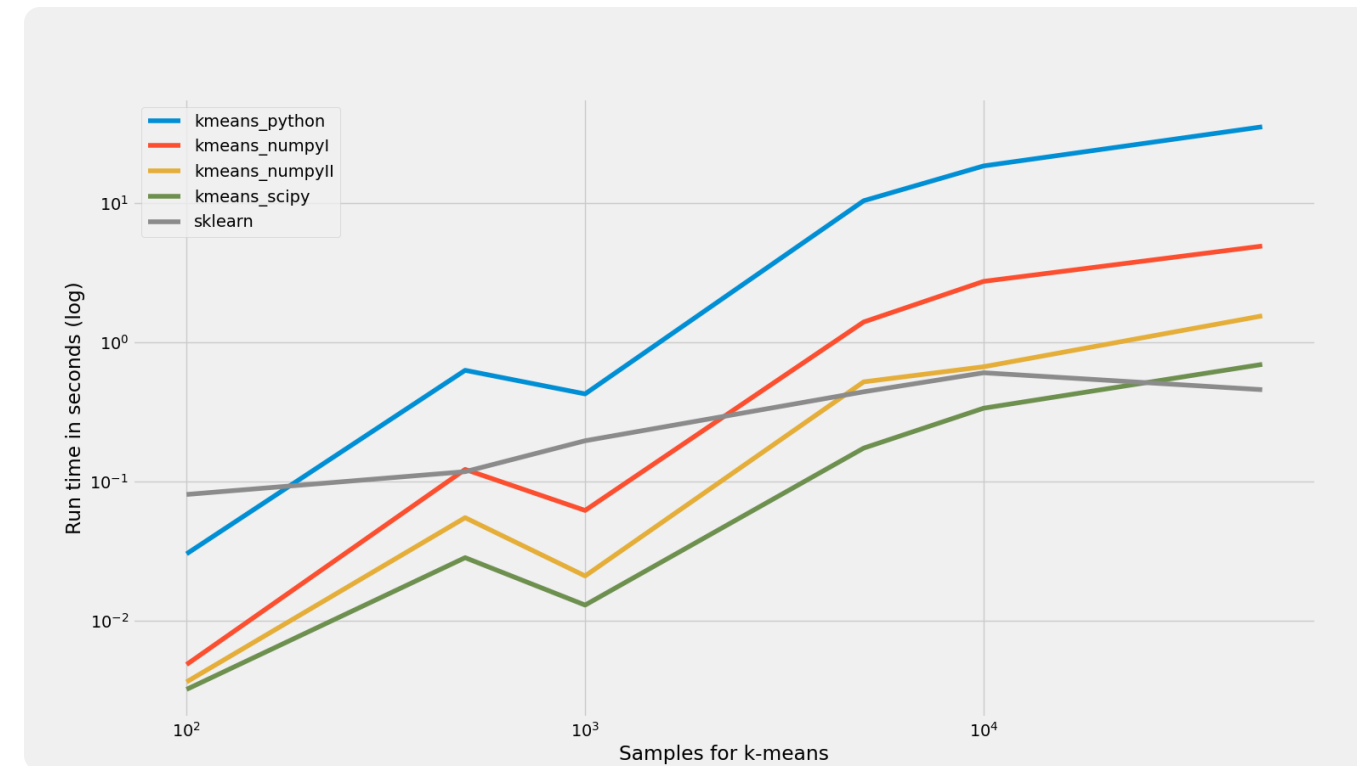
Method comparison



Method comparison



Method comparison



Note:

- the run time depends on the number of samples
- indeed, our implementations monotonically got better for the different variants
- the variant from the sklearn library is faster for large numbers of samples.

Other tools to improve speed (see excercises)

- Detailed guide on Python - C++ interfacing
- Cython python-like Code at C-like speeds; recommended for operations that cannot be done in NumPy, SciPy, etc.
- Numba for just-in-time compilation
- f2py call Fortran code

Questions?

Today's summary

Always profile your code, before you optimize it!

How to profile?

- We discussed profiling tools cProfile, timeit, line_profiler incl. visualization methods
- We learned about one of the core clustering algorithms. It's otherwise not important for this class, but I think it's algorithm you should know

How to optimize?

- Numpy is faster, always try to write vectorized code!
- Use *existing, optimized and specialized functions* (cdist, kmeans from a library, ...). This is also much better maintained
- For specialized functions, consider writing them in Fortran, C++ and link
- We learned about memoization and functools: lru_cache

After lunch:

- Monday 13 - 15: exercises working on v6 of your project
- Monday 15:15 - 16: my office hours at SV 2811