

CS 2112 Fall 2026

Assignment 2

Data Structures and Text Editing

Written problems due Thursday, September 24th, 11:59PM

Programming assignment due Tuesday, October 6th, 11:59PM

Text editors must store large dictionaries of words and quickly access them when performing common tasks such as word completion, spell checking, and text search. In this assignment you will implement core data structures and algorithms for a simplified text editor. The first part introduces a generic hash table and a prefix tree. The second part requires you to create plugins for a text editor that performs word completion and spell checking.

The last part contains written problems focusing on the concepts introduced in class.

This assignment will take some time. Get started early!

Updates

- 9/19 - Changed the value and next fields of `ListNode` in Written Problem 7.1 from private to package-private.

1 Instructions

1.1 Grading

Submissions will be graded for design, correctness, testing, and style. A good design makes the implementation easy to understand and maximizes code sharing. A correct program compiles without errors or warnings and behaves according to the requirements given here. A good test plan ensures good coverage of features and edge cases. A program with good style is clear, concise, and easy to read.

A few suggestions regarding good style may be helpful. You should use brief but mnemonic variable names and proper indentation. Your code should include comments as necessary to explain how it works, but without explaining things that are obvious.

1.2 Collaboration

You must work alone for this assignment. The course staff are happy to help with any difficulties that might arise. Use Ed for questions and don't be shy about coming to office hours if you need help. You are also welcome to discuss the assignment with other students, subject to the rules discussed on the website.

1.3 Documentation

For this assignment, we are especially looking for good documentation of the interfaces implemented by your data structures. Write Javadoc-compliant comments that crisply explain what all the methods do at a level of abstraction that enables a client to use your data structure effectively, while leaving out implementation details that a client does not need to know.

1.4 Restrictions

Your use of `java.util` will be restricted for this assignment. **Classes** from `java.util`, except for `Scanner`, may not be used anywhere in your code except in a JUnit test suite (see §5). The class `java.math.BigInteger` may not be used in your implementation either. **Interfaces** from `java.util` may be used anywhere in your code to guide your internal data structures.

While we require that you respect any interfaces we release, you are allowed (and even expected) to create your own classes and interfaces to solve portions of the assignment.

1.5 Generative AI

You are allowed to use Generative AI for this assignment **only** for writing test cases. However, you **must** still write some test cases by hand. Furthermore, in our experience, test cases written by GenAI are not up to the standards expected by this course. Remember that course policy requires that you document all prompts given to GenAI in `README.pdf`. Basic code completion in IntelliJ is still permitted. Consult with the course staff if you have specific questions.

1.6 Importing and Running

Download the release zip from CMSX and extract its contents to the place you keep your projects. In IntelliJ, select `File` → `Open` and select the unzipped folder. Ensure your project is using Java 21 by choosing `File` → `Project Structure` and checking to see that version 21 is selected under SDK.

Starting with this assignment, we will be using a system called Gradle in the release code. Gradle automatically adds any dependencies into your project without the need to add them manually. IntelliJ should automatically build the project using Gradle.

Once the build is done, you will have to set up the run configuration for the project. On the right sidebar of the IDE there will be an icon of an elephant that says Gradle when moused over, click on that. A sidebar will open, select `A2Release` → `Tasks` → `application` → and then double click `Run`. This will run the project and reveal the GUI you will be using. To stop running the project, close the GUI as you would any normal computer application. To rerun the application, you should now just be able to select the green play arrow at the top of the screen.

An alternative way to set up the run configuration is to click the drop-down menu to the left of the play button at the top of the screen and click `Edit Configurations`. This will open up the `Run/Debug Configurations` dialog. Now click the `+` on the top left of the screen and select Gradle. Then in the `Name` input box type something such as “Run A2”, then in the `Run` input field, simply type in “run”. Select `Apply`, then `OK`. You should now be able to click the green play arrow to run the application.

1.7 Tips

In this assignment, you will be modifying an application with a graphical user interface (GUI). The application has significant library dependencies because it builds on the JavaFX GUI library. To make sure you don’t run into headaches right before the deadline, start early to make sure that you have the right setup to successfully modify, compile, and run the application.

It will be tempting to think that you will be mostly done once you have implemented the data structures. But actually there are a few things to do after: profiling, performance optimization, integration into the GUI application. So aim to get the data structure part done well before the deadline.

2 Hash tables

Your task in this section is to implement a hash table with chaining. The [lecture notes](#) on hash tables have some helpful pointers, but we will also provide a high level overview here, since we won’t cover them for a few more lectures.

A hash table is a data structure which maintains key–value pairs. Each key is mapped to an index in an array using a hash function. Elements have a high probability of being hashed to unique indices, but in the case of a collision (multiple elements mapping to the same index) elements can either be stored in the same index through use of a linked list (chaining) or just stored in the next available index (probing).

The benefits of a hash table are that common data structure operations have a significantly better run time in the average case. For example, lookup in an array is $O(n)$ but for a hash table, it is $O(1)$. You will learn more about this in lecture, but getting a head start and understanding it on a high level can help with this assignment.

2.1 Collisions

You should use chaining to handle collisions. You are expected to keep track of the load factor and to resize your table whenever the load factor crosses a threshold. A smart choice of load factor will keep memory usage reasonable while avoiding collisions. You are only required to resize your hash table when the load factor goes **above** your chosen threshold. In other words, resizing down is not required.

2.2 Implementation

You will be implementing the class `HashTable<K,V>`. Your hash table should implement the interface `java.util.Map<K,V>`, which is generic. The methods `containsKey`, `get`, `put`, and `remove` should have expected $O(1)$ (constant) running time. Your hash table should take up $O(n)$ (linear) space, where n is the number of entries in the hash table.

The `HashTable<K,V>` constructor takes an integer parameter `numElements` which should be used as a hint to determine the starting number of buckets within the constructed object.

The method `keySet()` should return an object implementing `java.util.Set<K>`. This object must support the following methods:

- `size()` and `isEmpty()`;
- `toArray()` and `contains(Object)`.

The `toArray(T[])` method and the remaining methods may throw `UnsupportedOperationException`. The object returned by this method must be updated automatically as the hash table is modified.

The method `hashCode()`, which is defined for every Java object, can be used by a hash function that you create to compute the bucket in which to place each object. However, since `hashCode()` is not required to produce results that behave as if they are random, you don't want to use `hashCode()` directly to compute the bucket index. For example, the default implementation of `hashCode()` returns the object's memory address, therefore only produces numbers that are multiples of 4. Another hash function is needed to provide diffusion throughout the buckets. The class `java.security.MessageDigest` provides high-quality hash functions that can be used for this purpose, although they are more expensive than necessary for most applications. The course notes have tips on how to design a `hashCode()` method; see also this [Wikipedia page](#).

3 Prefix trees

A prefix tree, also known as a **trie**,¹ is a data structure tailored for storing and retrieving strings. The root node represents the empty string.² Each possible next character branches to a different child node. Strings stored in the trie must be inserted explicitly by the user; prefixes of such strings, although they occur along paths in the trie, are not considered to be stored in the trie unless they have been explicitly inserted.

For example, the trie of Fig. 1 contains the four strings `COW`, `CS`, `CS2110`, and `CS2112`. The strings `C`, `CS211`, `C0`, and the empty string, although they appear as prefixes of strings stored in the trie, are not considered to be stored in the trie themselves.

If a string is stored in the trie, there is a unique node corresponding to that string and a unique path from the root down to that node obtained by tracing the characters in the string. That node can contain a boolean flag to indicate that that string has been stored in the trie. There is no need to store the string itself

¹Pronounced like "try".

²Note that the empty string is "", the string of length 0, not null.

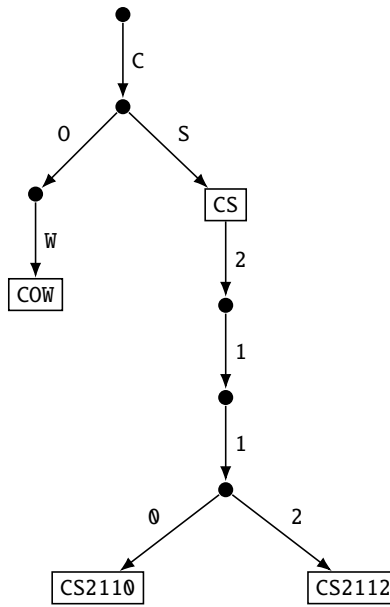


Figure 1: A trie containing the strings COW, CS, CS2110, and CS2112.

at that node; the string can be recovered by tracing the path from the root down to that node, keeping track of the characters along the way.

3.1 Implementation

Implement the provided `Trie` class. The operations `insert`, `delete`, and `contains` should have $O(k)$ running time, where k is the length of the string. In other words, the running time of these operations should be proportional to the length of the given string. Your trie should also implement the method `closestWordToPrefix()`, which returns the shortest entry in the trie having the given prefix. This shortest string can be found using breadth-first search.

The method `closestWordToPrefix()` should be case-sensitive. For example, it should report CS2110 or CS2112 if the argument is CS211, but not if the argument is cs211.

Your trie should be able to handle strings containing any character, not just alphanumeric characters.

4 Text editor

The text editor supports text search, spell checking, and autocompletion. These features are specified by the interfaces `SearchModule`, `SpellCheckModule`, and `AutoCompleteModule`. You are to provide implementations. The factory class `ModuleFactory` contains factory methods that should access your implementations. Instances returned from the factory methods are used by the main text editor program.

4.1 Architecture

The text editor project is broken up into three packages. The `editor` package contains the code for the user interface. The `modules` package contains all of the plugins providing functionality for text search, spell checking, and autocompletion. The `util` package contains all of the data structures you will implement. These data structures store and manipulate data for the plugins. While all the code you are required to write resides in the `modules` and `util` packages, you are welcome to look inside the `editor` package to get a taste of graphical user interface (GUI) code.

4.2 Dictionary file

After the text editor is started, spell checking and autocompletion are unavailable until a dictionary file is loaded. Any newline-separated list of words will work as a dictionary file. WinEdt provides [such a file](#). On Macintosh and most Linux distributions, a good dictionary file can be found at `/usr/share/dict/words`. To load a dictionary file, click the top left button of the text editor.

4.3 User interaction

If your modules work correctly, word-completion suggestions from the autocomplete module should be displayed in the lower-left corner of the editor window. Misspelled words should be highlighted if you click the “check” button in the top left. To reset spell checking, click the adjacent “X” button. Additionally, the time spent spell checking should be reported in the lower-right corner after each run of spell checking. If you enter a string in the search window at the bottom and click the search button, the first occurrence of this string should be highlighted.

4.4 Implementation

You should not modify any code in the `editor` package. The functionality for the editor will come from your implementations of the interfaces in the `modules` package. Your implementation of these interfaces should stand alone and follow the given specifications without modifications to the `editor` package.

Spell checking and autocompletion should both convert dictionary words to lowercase before being entered. The text editor will automatically convert text from the GUI to lowercase before passing it into `getWordForPrefix` or `isValidWord`. It does not do this for search queries, which should be case-sensitive.

5 Testing

In addition to the code you write for the data structures and text editor plugins, you should also submit any tests that you write. Testing is a component of the grade for this assignment.

You should implement your test cases using JUnit, a framework for writing test suites. IntelliJ makes running JUnit tests very easy; just click the green arrow next to the test class name to run all tests, or run individual test methods by clicking the green arrow next to the one you would like to run.

You should not only test whether the program works correctly from the command line interface, but also write test cases for each of the data structures you implement.

Test cases should be placed in a top-level directory named `src/test`, whereas the rest of your implementation would be in `src/main`.

There are several good strategies for writing test cases. In **black-box functional testing**, the tester defines input-output pairs in which the inputs provide good coverage of the input space. Each input is accompanied by the expected output as defined by the specification. We expect you to define functional test cases for your program as a whole and for each data structure you implement.

A second approach to testing is **random testing**, in which the inputs are generated randomly but in a way that satisfies the preconditions. A random test case might generate a sequence of randomly chosen inputs to a single method or to a randomly chosen method from a set of methods. This form of testing can catch bugs simply when the code fails with an exception or assertion error. Often an effective way to randomly test functional correctness is to test whether the behavior of the code matches that of a simple **reference implementation** on which the same operations are performed. For example, the `java.util` libraries may be used to build simple reference implementations for each of the abstractions you are implementing. We expect you to use random testing on at least one abstraction you develop in this assignment.

6 Performance and Correctness

Both correctness and performance are important when we evaluate how well the editor plugins work.

6.1 Performance

Performance analysis is a component of the grade for this assignment. You should choose data structure(s) wisely to be efficient in both memory usage and run time. Justify your design in `README.pdf`.

In addition to providing a qualitative justification, you should conduct performance tests on your hash table implementation by inserting random unique strings (the details of how these are generated are up to you). For each of the following measurements, create a graph plotting it against the number of elements in your hash table.

- Put time
- Get time
- Number of empty buckets
- Number of collisions

Hash table size should range from 0 to at least 100,000 elements, with data points for each number of elements. Additionally, each data point should be the average of 5 timed method calls, (i.e. you should run 5 trials, each on a different hash table instance, and average them to create your data points). Method calls should be timed using `System.nanoTime()`.

Include a line of best fit on each graph. Excel or Google Sheets can be helpful in creating these.

Include these graphs in the file `perf.pdf`, with a brief explanation of why your results show the following properties of a good hash table implementation:

- The put and get methods are $O(1)$.
- The hash function produces reasonable diffusion.

Be aware that Java programs run very slowly when they first start, because libraries are being loaded and code is run in a slower, interpreted mode initially. Frequently used code is compiled “just in time” by the JIT compiler to machine code that runs at least an order of magnitude faster. Try to collect performance measurements only after the code being measured has run for 10 seconds.

6.1.1 Profiling

In lab, we covered profiling with VisualVM/IntelliJ, which can give a lot of insight about where time is being spent in your code. You are encouraged to use profiling throughout your development process, but you are required to profile your final product using VisualVM/IntelliJ, and include a few paragraphs in `README.pdf` describing your results and what they mean about the performance of your code.

Your response should include answers to the following questions:

- According to your CPU profiler results, which operation or code path in your hash table implementation (other than put or get) consumes the most CPU time? Why do you think this is the case?
- How does resizing (rehashing) affect both CPU time and memory usage in your hash table? Approximately what fraction of total runtime and memory allocation does resizing account for in your tests?
- According to the memory profiler, which data structure or operation in your Trie implementation uses the most memory? Explain why this occurs and whether it matches your expectations.
- Did you observe any unexpected CPU or memory performance bottlenecks in your data structures? If so, what were they and how might you address them?

6.2 Correctness

A good way to see if your tests are actually *testing* your code well is to try and trace what branches of code are executed. For instance, you may have inadvertently constructed a test suite that tests one method very thoroughly, but that omits another method altogether. You also may be always avoiding one buggy `else if` statement that only executes for edge cases and all your tests pass because they bypass the bugs. Regardless of whether you choose to approach testing from a randomized, glass or black box method, you should always strive to make sure your code runs at least once in a test suite for sanity's sake.

Tracing these branches manually can be rather difficult, but there are tools that can help you design your tests to achieve better overall **coverage** of your code. IDEs like IntelliJ often have **built in** coverage tools. These tools helpfully tell you exactly how much of your code is being executed in your unit tests. You should use Run with Coverage to achieve as close to 100% coverage as possible on your tests for `HashTable`.

To run a specific test with coverage, you should already have an existing active test/run configuration that you wish to run with coverage. Then, you can select `Run → Run <configuration> with Coverage` from the menu bar; it should look like a run with a shield icon. If the configuration runs as expected and no coverage pops up, click into the edit run configuration window, and scroll all the way to the bottom to add the specific directories that you are interested in tracing coverage for. Once you have successfully run a test with coverage on, you should see a Coverage tab pop up. If not, you can go to `View → Tool Windows → Coverage` to open it. The fourth button on the Coverage tool window will allow you to export your test results as an HTML file; include this export in the Coverage/ subdirectory with your final submission in your ZIP.

7 Written problems

7.1 Abstraction

The standard Java interface `SortedSet` describes a set whose elements have an ordering. Abstractly, the set keeps its elements in sorted order. Here is a much simplified version:

```
1  /// A set of unique elements kept sorted in ascending order according
2  /// to some ordering.
3  interface SortedSet<T> {
4      /// Effect: Add x to the set if it is not already there.
5      void add(T x);
6
7      /// Whether x is in the set.
8      boolean contains(T x);
9
10     /// Effect: Remove element x if present.
11     void remove(T x);
12
13     /// The first (least) element in the set.
14     T first();
15 }
```

1. The specifications of some of these methods are incomplete. Clearly identify the problems and write better specifications for the methods that need to be improved. You may change method signatures if you justify the change.
2. There are many ways to implement this set abstraction. One possibility is as a linked list data structure in which there are no duplicates and the elements are kept in sorted order:

```
public class SortedList<T> implements SortedSet<T> {
    /// Representation: A linked list of values starting at 'head', which may be 'null'
    /// to represent an empty list.
    ///
    /// <p>Invariant: the list nodes starting from 'head' have values in ascending
```

```

    /// sorted order according to 'comparator', with no duplicates.
    /// This is the order in which they occur in the set.
    ///
    private @Nullable ListNode<T> head;
    private Comparator<T> comparator;
}

public class ListNode<T> {
    T value;
    ListNode<T> next;

    public ListNode(T v, ListNode<T> n) {
        value = v;
        next = n;
    }
}

```

The SortedList implementation is obviously incomplete. Give the most efficient, concise code you can to implement the first and remove methods, taking into account the representation and class invariant.

- Now, suppose we want a different implementation UnsortedList that is similar to SortedList and uses the same ListNode class, but has a much weaker class invariant:

```

public class UnsortedList<T> implements SortedSet<T> {
    /// Representation: the set contains the values contain
    /// in the linked list of values starting at 'head'. The
    /// order of the values in the list is arbitrary, but
    /// the order of the values in the set is determined by 'comparator'.
    private @Nullable ListNode<T> head;
    private Comparator<T> comparator;
}

```

UnsortedList should still correctly implement the SortedSet interface. Implement the add, first, and remove methods as simply and concisely as you can, taking into account the representation and class invariant.

Since SortedList and UnsortedList implement the same specification, the client should not be able to tell which one is being used, except perhaps by timing.

- Briefly discuss the advantages and disadvantages of each of these two implementations. Under what conditions it would be more appropriate to use SortedList? ...UnsortedList?

7.2 Asymptotic complexity

Recall that a function $f(n)$ is $O(g(n))$ if there exist positive constants k and n_0 such that for all $n \geq n_0$, $f(n) \leq kg(n)$. The constants k and n_0 together are a **witness** to the fact that $f(n)$ is $O(g(n))$.

- Consider the code snippet below. Give a tight bound on its time complexity using big-O notation, and briefly justify your answer.

```

1 for (int i = 5; i < n; i++) {
2     if (i % 2 == 0) {
3         for (int j = i + 1; j < n; j++) {
4             for (int k = 7; k < 70000; k++) {
5                 System.out.println("2112_is_great!");
6             }
7         }
8     }
9 }

```

- Prove that $n^2 \lg n$ is $O(n^3)$. Be sure to specify a witness pair (k, n_0) .
- Prove that if $f_1(n)$ and $f_2(n)$ are both $O(n^2)$, then $f_1(n) + f_2(n)$ is $O(n^2)$.

7.3 Hashing

8. Show the state of the underlying array of a hash table, when implemented with chaining and then with linear probing using a stride of 1. Assume the hash function is simply n modulo the length of the array. The elements inserted into the array are 4, 15, 54, 43, 25, 42, 30, 2112, 2025, 21, 9, 65, 44, 219.

The initial length of the array is 5, and the maximum load factor for the chaining implementation is 2, and for the probing implementation is 1. The array size is doubled when the maximum load factor is reached. Assume that elements are rehashed in the order in which they were inserted.

7.4 Recursion

9. The iterative function `suffixEquals` checks whether the suffixes of two strings are equal without using `String.substring`. Convert it into a tail-recursive function.

```
1 /**
2  * Checks if the suffix of {@code str1} starting at {@code idx1}
3  * is equal to the suffix of {@code str2} starting at {@code idx2}.
4  *
5  * Requires: 0 <= idx1 < str1.length() and
6  *           0 <= idx2 < str2.length().
7  */
8 boolean suffixEquals(String str1, String str2, int idx1, int idx2) {
9     while (idx1 < str1.length() && idx2 < str2.length()) {
10         if (str1.charAt(idx1) != str2.charAt(idx2)) {
11             return false;
12         }
13         idx1++;
14         idx2++;
15     }
16     return idx1 == str1.length() && idx2 == str2.length();
17 }
```

10. Consider the following tail-recursive function `gcd`, which computes the greatest common divisor of two numbers using Euclid's algorithm. Convert it into an iterative function.

```
1 /**
2  * Computes the greatest common divisor of {@code a} and {@code b}.
3  *
4  * Requires: a >= 0 and b >= 0.
5  */
6 int gcd(int a, int b) {
7     return b == 0 ? a : gcd(b, a % b);
8 }
```

11. Complete the following function `makeMonotonic` using a strictly recursive $O(n)$ solution, where n is the number of nodes in the list. Do not use loops. Use `compareTo()` to determine the ordering of node values.

For example, the input list 5 -> 2 -> 6 -> 3 -> 3 -> 1 should become 6 -> 3 -> 1.

```
1 /**
2  * A linked list class. A node n is the end of the list if
3  * n.next == null.
4  */
5 class ListNode<T> extends Comparable<T> {
6     T value;
7     ListNode<T> next;
8 }
9
10 /**
11  * Makes a linked list monotonic by deleting all nodes for which
12  * there is a subsequent node greater than or equal to it.
13  */
```

```

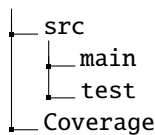
14 void makeMonotonic(ListNode<T> head) {
15     // Your code goes here.
16 }

```

8 Submission

You should submit these items on CMS:

- **README.pdf:** This file should contain your name, your netID, all known issues with your submitted code, the names of anyone you discussed the assignment with (including clarifications from course staff), and any other sources that should be acknowledged.
In addition, you should briefly describe your design, noting any interesting design decisions you encountered, and briefly discuss your testing strategy and profiling results. You can follow the [design overview guidelines](#) on the course web site.
- **written.txt or written.pdf:** This file should include your response to the written problems.
- **perf.pdf:** This file should include your performance analysis.
- **Source code:** Please compress your code into a zip file with the following structure:



Because this assignment is more open than the last, you should include all source code required to compile and run your project. All source code should reside in the `src` directory with an appropriate package structure. You should include code for all your test cases in `test`. Subpackages are permitted. Do not include any `.class` files or any other extraneous files. Please export your code coverage report into a subdirectory named `Coverage` in the root directory of your zip.

All `.java` files should compile and conform to the prototypes we gave you. We write our own classes that use your classes' public methods to test your code. *Even if you do not use a method we require, you should still implement it for our use.*

9 Optional: Probabilistic Data Structures

This is not an official part of the assignment. No extra credit will be given, but you are welcome to give it a try just for fun and good karma.

You may incorporate bloom filters into your text editor application as you see fit, but be careful not to break any required functionality while doing so. Document anything you do that goes beyond what is mandatory in your design overview.

9.1 Bloom Filters

A Bloom filter is a probabilistic constant-space data structure for maintaining a set of elements and testing whether a given element is in the set. It is probabilistic in the sense that false positives may occur with small probability (that is, an element may be reported to be in the set when it is not), but false negatives never occur (that is, if an element is reported not to be in the set, then it is definitely not in the set).

An empty Bloom filter is a bit array of 0s. To insert an element into a Bloom filter, put the element through k different hash functions. Use the results of these hash functions as indices into the bit array. Set those k bits in the bit array to 1.

To determine if an element is in the Bloom filter, check all of its hash indices. If all of them are 1 in the bit array, report that the element is in the set. If at least one of them is 0, report that the element is not in the set.

If the objects contained in the Bloom filter are strings, the k different hash functions can be simulated with a single hash function by appending a different single character (e.g., a, b, c, ...) to the end of the string before hashing.

9.1.1 Example of a false positive

Consider a Bloom filter for strings represented by a bit array of length 2, initially empty. Suppose only one hash function is used to index strings. First, the string CS2112, whose (hypothetical) hash value is 0, was inserted into the Bloom filter, setting the 0th bit to 1 in the bit array. Now, to check whether CS2110, whose hypothetical hash value is also 0, is in the Bloom filter, we check if the bit at position 0 is 1. Since this is the case, we conclude that the Bloom filter does contain the String CS2110 when in fact it does not.

A larger bit array, more hash functions, and better quality hash functions all reduce the likelihood of false positives.

9.2 Implementation

We have provided a `BloomFilter` class for you to implement if you choose to complete this optional part of the assignment.