

МИНИСТЕРСТВО НАУКИ И ВЫСШЕГО ОБРАЗОВАНИЯ
РОССИЙСКОЙ ФЕДЕРАЦИИ
федеральное государственное автономное образовательное учреждение
высшего образования «Санкт-Петербургский политехнический
университет Петра Великого»

Институт компьютерных наук и кибербезопасности

Высшая школа технологий искусственного интеллекта

Направление: 02.03.01 «Математика и компьютерные науки»

Лабораторная работа №5
«Реализация LL(1)-анализатора»
по дисциплине
«Математическая логика и теория автоматов»
Вариант 15

Студент,

группы 5130201/20102

_____ Тищенко А. А.

Преподаватель

_____ Востров А. В.

«_____» _____ 2025г.

Санкт-Петербург, 2025

Содержание

Введение	4
1 Математическое описание	5
1.1 Анализ исходной грамматики и её модификация	5
1.2 Иерархия грамматик Хомского	5
1.3 Контекстно-свободные грамматики	6
1.4 LL(k)-анализаторы и LL(k)-грамматики	6
1.5 Алгоритм построения LL(1) таблицы синтаксического анализа	8
1.6 Множества FIRST и FOLLOW для заданной грамматики	9
1.7 Таблица синтаксического анализа для заданной грамматики	9
1.8 Алгоритм разбора предложений	10
1.9 Алгоритм генерации цепочек	10
1.10 Семантические действия	11
2 Особенности реализации	13
2.1 Задание правил грамматики	13
2.2 Класс <code>Grammar</code>	13
2.2.1 Метод <code>_parse_productions</code>	14
2.2.2 Метод <code>_calculate_first_sets</code>	14
2.2.3 Метод <code>_calculate_follow_sets</code>	15
2.2.4 Метод <code>_first_of_sequence</code>	17
2.2.5 Метод <code>_fill_lookup_table</code>	17
2.2.6 Метод <code>format_rules</code>	18
2.2.7 Метод <code>format_lookup_table</code>	19
2.2.8 Метод <code>format_first_sets</code>	19
2.2.9 Метод <code>format_follow_sets</code>	20
2.2.10 Метод <code>analyze</code>	20
2.2.11 Метод <code>generate</code>	22
2.2.12 Метод <code>generate_derivation_steps</code>	22
2.3 Функция <code>load_grammar</code>	23
2.4 Функция <code>main</code>	24
3 Результаты работы программы	26
Заключение	28

Введение

Лабораторная №5 заключается в следующем. Для заданной вариантом грамматики необходимо:

- Построить множества FIRST и FOLLOW для каждого нетерминала.
- Построить таблицу выбора (lookup table).
- Реализовать детерминированный левый анализатор (проверка принадлежности цепочки грамматике).
- Назначить семантические действия части заданных продукций (например, генерация машинноориентированного кода).

Исходная грамматика в варианте 15 состояла из следующих продукций:

```
S -> B A b
A -> a A B C | b B | a
B -> b
C -> c A
```

Однако, как будет показано в разделе математического описания, данная грамматика не является LL(1)-грамматикой из-за неоднозначности в выборе продукций для нетерминала A . Поэтому в работе используется модифицированная версия грамматики с удаленной продукцией $A \rightarrow a$.

1 Математическое описание

1.1 Анализ исходной грамматики и её модификация

Исходная грамматика варианта 15 содержала следующие productions:

```
S -> B A b
A -> a A B C | b B | a
B -> b
C -> c A
```

При попытке построения LL(1)-анализатора для данной грамматики обнаруживается неоднозначность в таблице синтаксического анализа. Конкретно, при рассмотрении нетерминала A и входного терминала a возникает конфликт между двумя productions:

- $A \rightarrow a$
- $A \rightarrow aABC$

Обе productions начинаются с терминала a , что означает, что $a \in \text{FIRST}(a)$ и $a \in \text{FIRST}(aABC)$. Следовательно, анализатор не может однозначно определить, какую production применить при встрече символа a на входе и нетерминала A на вершине стека.

Для получения корректной LL(1)-грамматики была выполнена модификация: удалена production $A \rightarrow a$. Таким образом, итоговая грамматика, используемая в данной работе, имеет вид:

```
S -> B A b
A -> a A B C | b B
B -> b
C -> c A
```

Данная модификация устраняет неоднозначность и позволяет построить корректную таблицу синтаксического анализа для LL(1)-анализатора. Язык, порождаемый модифицированной грамматикой, остается достаточно выразительным для демонстрации принципов работы LL(1)-анализатора.

1.2 Иерархия грамматик Хомского

Иерархия Хомского — классификация формальных языков и формальных грамматик, согласно которой они делятся на 4 типа по их условной сложности. На Рис. 1 представлены все четыре типа грамматик (по типу productions, из которых они состоят) и отношения между ними.

Любая автоматная грамматика является контекстно-свободной, любая контекстно-свободная грамматика является контекстно-зависимой, любая контекстно-зависимая грамматика является неограниченной.

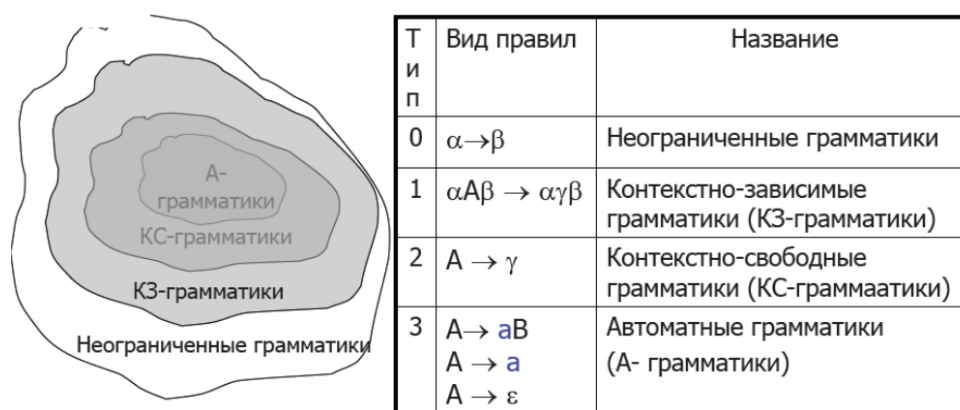


Рис. 1. Иерархия грамматик Хомского

1.3 Контекстно-свободные грамматики

Контекстно-свободные грамматики являются частным случаем формальных грамматик, в которых левые части всех продукций содержат только один нетерминальный символ. Это означает, что все продукции имеют вид $A \rightarrow \beta$, где A — нетерминальный символ, а β — произвольная цепочка терминалов и нетерминалов.

Модифицированная грамматика этого варианта лабораторной работы задаётся в виде следующего списка продукций:

```

S -> B A b
A -> a A B C | b B
B -> b
C -> c A

```

Она, очевидно, является контекстно-свободной, так как в левых частях всех продукций присутствует только один нетерминальный символ. Также можно сразу заметить, что грамматика не является автоматной, так как в списке есть продукции, в правых частях которых присутствуют несколько нетерминалов.

Язык, порождаемый грамматикой, очевидно, не сложнее контекстно-свободного. Также, по продукции $A \rightarrow a A B C$, можно сделать вывод, что язык не является автоматным. Эта продукция является рекурсивной и порождает вложенную структуру, которая не может быть описана конечным автоматом.

Пример дерева вывода для цепочки языка, порождаемого грамматикой, представлен на Рис. 2.

1.4 LL(k)-анализаторы и LL(k)-грамматики

Синтаксический LL-анализатор — это нисходящий синтаксический анализатор для некоторого подмножества контекстно-свободных грамматик, известных как LL-грамматики. При этом не все контекстно-свободные грамматики являются LL-грамматиками.

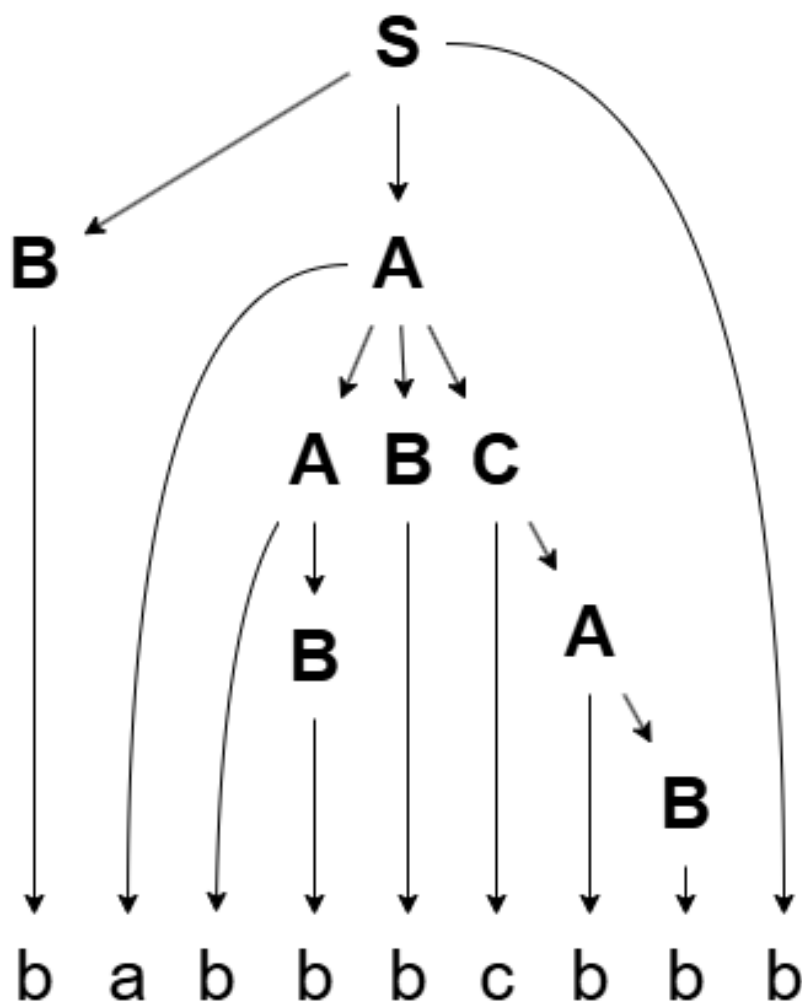


Рис. 2. Дерево вывода для цепочки *babbbcbbb*.

Буквы L в выражении «LL-анализатор» означают, что входная строка анализируется слева направо (left to right), и при этом строится её левосторонний вывод (leftmost derivation).

LL-анализатор называется LL(k)-анализатором, если данный анализатор использует предпросмотр на k токенов (лексем) при разборе входного потока. Грамматика, которая может быть распознана LL(k)-анализатором без возвратов к предыдущим символам, называется LL(k)-грамматикой. Язык, который может быть представлен в виде LL(k)-грамматики, называется LL(k)-языком.

Как будет показано далее, грамматика, рассматриваемая в этом варианте лабораторной работы является LL(1)-грамматикой, так как для неё можно однозначно заполнить таблицу LL(1) синтаксического анализа. Это означает, что для распознавания любой цепочки грамматики, анализатору достаточно просматривать входной поток только на один символ вперёд для принятия решения о том, какое правило грамматики необходимо применить.

1.5 Алгоритм построения LL(1) таблицы синтаксического анализа

Чтобы заполнить таблицу синтаксического анализа, необходимо установить, какое правило грамматики синтаксический анализатор должен выбрать, если нетерминальное A находится на вершине его стека и символ a — во входном потоке. Легко заметить, что такое правило должно иметь форму $A \rightarrow w$ и что у языка, соответствующего w , должна быть по крайней мере одна строка, начинающаяся с a .

С этой целью мы определяем **множество FIRST()** для w , обозначенное как $\text{FIRST}(w)$, как множество терминалов, которые могут быть найдены в начале любой строки в w , и ε , если пустая строка также принадлежит w .

Учитывая грамматику с правилами $A_1 \rightarrow w_1, \dots, A_n \rightarrow w_n$, можно вычислить $\text{FIRST}(w_i)$ и $\text{FIRST}(A_i)$ для каждого правила следующим образом:

1. Инициализировать каждое множество $\text{FIRST}(A_i)$ пустым множеством.
2. Добавить $\text{FIRST}(w_i)$ к $\text{FIRST}(A_i)$ для каждого правила $A_i \rightarrow w_i$, где $\text{FIRST}(w_i)$ определяется следующим образом:
 - $\text{FIRST}(aw') = \{a\}$ для каждого терминала a
 - $\text{FIRST}(Aw') = \text{FIRST}(A)$ для каждого нетерминального A с $\varepsilon \notin \text{FIRST}(A)$
 - $\text{FIRST}(Aw') = (\text{FIRST}(A) \setminus \{\varepsilon\}) \cup \text{FIRST}(w')$ для каждого нетерминального A с $\varepsilon \in \text{FIRST}(A)$, включая случай $A_i \rightarrow A$, то есть $w' = \varepsilon$, в этом случае $\text{FIRST}(Aw') = \text{FIRST}(A)$
 - $\text{FIRST}(\varepsilon) = \{\varepsilon\}$
3. Повторять шаг 2, пока в множествах **FIRST** происходят изменения.

Однако множеств **FIRST()** недостаточно, чтобы построить таблицу синтаксического анализа. Так происходит, потому что правая сторона w правила могла бы в конечном счёте быть приведена к пустой строке. Таким образом синтаксический анализатор должен также использовать правило $A \rightarrow w$, если $\varepsilon \in \text{FIRST}(w)$ и во входном потоке символ, который может следовать за A . Поэтому также необходимо построить **множество FOLLOW()** для A , которое определяется как множество терминалов a , таких что существует строка символов $\alpha A a \beta$, которая может быть получена из начального символа.

Вычисление множеств **FOLLOW()** для нетерминалов в грамматике выполняется следующим образом:

1. Инициализировать $\text{FOLLOW}(S) = \{\$ \}$, а все остальные множества $\text{FOLLOW}(A_i)$ пустыми, где $\$$ — специальный терминал, служащий для указания конца входной цепочки.
2. Если есть правило формы $A_j \rightarrow w A_i w'$, тогда:
 - Если терминал $a \in \text{FIRST}(w')$, то добавить a в $\text{FOLLOW}(A_i)$

- Если $\varepsilon \in \text{FIRST}(w')$, то добавить $\text{FOLLOW}(A_j)$ в $\text{FOLLOW}(A_i)$
- Если w' имеет длину 0, то добавить $\text{FOLLOW}(A_j)$ в $\text{FOLLOW}(A_i)$

3. Повторять шаг 2, пока в множествах $\text{FOLLOW}()$ происходят изменения.

Теперь можно точно определить, какие правила будут содержаться в таблице синтаксического анализа. Если $T[A, a]$ обозначает ячейку в таблице для нетерминального A и терминала a , то:

- $T[A, a]$ содержит правило $A \rightarrow w$ тогда и только тогда, когда:
 - $a \in \text{FIRST}(w)$ при проходе правила $A \rightarrow w$, или
 - $\varepsilon \in \text{FIRST}(w)$ и $a \in \text{FOLLOW}(A)$

Если таблица будет содержать не более одного правила в каждой ячейке, то синтаксический анализатор сможет однозначно определить, какое правило необходимо применить на каждом шаге разбора. В этом случае грамматику является **LL(1)** грамматикой.

1.6 Множества FIRST и FOLLOW для заданной грамматики

Приведём множества FIRST и FOLLOW для всех нетерминалов заданной грамматики.

Множества FIRST:

$\text{FIRST}(S) = \{b\}$
 $\text{FIRST}(A) = \{a, b\}$
 $\text{FIRST}(B) = \{b\}$
 $\text{FIRST}(C) = \{c\}$

Множества FOLLOW:

$\text{FOLLOW}(S) = \{\$ \}$
 $\text{FOLLOW}(A) = \{b\}$
 $\text{FOLLOW}(B) = \{a, b, c\}$
 $\text{FOLLOW}(C) = \{b\}$

1.7 Таблица синтаксического анализа для заданной грамматики

Таблица синтаксического анализа для заданной грамматики представлена в таблице 1.

Таблица 1. Таблица синтаксического анализа для заданной грамматики.

	a	b	c	\$
S	-	B A b	-	-
A	a A B C	b B	-	-
B	-	b	-	-
C	-	-	c A	-

Если при распознавании цепочки синтаксический анализатор встречается ячейку с прочерком, то это означает, что входная цепочка не принадлежит языку, порождаемому грамматикой. В таком случае синтаксический анализатор завершает работу и возвращает сообщение об ошибке.

1.8 Алгоритм разбора предложений

Перед запуском алгоритма разбора в стек помещаются два символа:

- Специальный символ \$ — признак конца входной цепочки.
- Начальный символ грамматики (на вершину стека).

Синтаксический анализатор выполняет три различных вида действий в зависимости от того, находится ли на вершине стека нетерминал, терминал или специальный символ \$.

1. **Если на вершине стека — нетерминал**, то в таблице синтаксического анализа ищется правило грамматики, находящееся на пересечении строки и столбца, соответствующих этому нетерминалу на вершине стека и текущему входному символу. Если же в указанной ячейке таблицы правило отсутствует — синтаксический анализатор останавливается и сообщает об ошибке.
2. **Если на вершине стека — терминал**, то синтаксический анализатор сравнивает его с текущим входным символом. Если они равны, то входной символ считывается, а соответствующий символ из вершины стека — удаляется.
3. **Если на вершине стека — специальный символ \$**, и текущий символ на ленте также равен \$, то синтаксический анализатор сообщает об успешном распознавании цепочки. В противном случае — сообщает об ошибке. В обоих случаях происходит останов анализатора.

Эти шаги повторяются до тех пор, пока не произойдёт останов. После останова мы получаем либо сообщение об ошибке, либо сообщение об успешном распознавании цепочки.

1.9 Алгоритм генерации цепочек

Генерация цепочек происходит следующим образом:

1. В список помещается начальный символ грамматики.
2. Осуществляется проход по списку. Первый встреченный нетерминал заменяется на его случайно выбранную продукцию, после чего проход начинается заново.
3. Алгоритм завершается, когда в списке не осталось нетерминалов.

Этот алгоритм не позволяет контролировать длину предложения. Также алгоритм может не сходиться в общем случае. Но вероятность того, что алгоритм попадет в петлю убывает геометрически с ростом длины цепочки, так что на практике этот метод работает.

1.10 Семантические действия

Каждой продукции можно сопоставить некоторое семантическое действие, которое будет выполняться при применении этой продукции, во время распознавания цепочки. Таким образом, можно, например, генерировать машинный код.

В данной лабораторной работе в качестве семантических действий решено было использовать генерацию инструкций для Java Virtual Machine. JVM это стековая виртуальная машина, её инструкции изменяют содержимое стека. Примеры инструкций:

- `iconst_1` — положить целочисленное значение 1 на вершину стека.
- `isub` — вычесть верхнее целочисленное значение стека из предыдущего и положить результат на вершину стека.
- `iadd` — сложить два верхних целочисленных значения стека и положить результат на вершину стека.

При распознавании цепочки можно, например, генерировать инструкции, в результате выполнения которых на вершине стека окажется следующее целочисленное значение:

$$\#b - 2 \times \#a + 3 \times \#c$$

где $\#a$, $\#b$, $\#c$ — это количества вхождений символов a , b , c в цепочку соответственно.

Для этого продукции были сопоставлены со следующими инструкциями:

- $S \rightarrow BAb \rightsquigarrow \text{iconst_1}$
- $A \rightarrow aABC \rightsquigarrow \text{iconst_2 isub}$
- $A \rightarrow bB \rightsquigarrow \text{iconst_1 iadd}$
- $B \rightarrow b \rightsquigarrow \text{iconst_3 iadd}$
- $C \rightarrow cA \rightsquigarrow \text{iconst_3 iadd}$

Тогда при левом выводе цепочки **babbbcbbb**, будет сгенерирована следующая последовательность инструкций:

```
iconst_1      # [1] - значение на вершине стека
iconst_1 iadd # [2]
iconst_2 isub # [0]
iconst_1 iadd # [1]
iconst_1 iadd # [2]
iconst_1 iadd # [3]
iconst_3 iadd # [6]
iconst_1 iadd # [7]
iconst_1 iadd # [8]
```

Значение на вершине стека в конце выполнения всех инструкций равно 8. Поскольку в цепочке **babbbcbbb** 7 символов *b*, 1 символ *a* и 1 символ *c*, то и значение выражения $\#b - 2 \times \#a + 3 \times \#c$ для этой цепочки равно 8.

2 Особенности реализации

2.1 Задание правил грамматики

Грамматика задаётся в виде текстового файла, в котором каждая строка соответствует одной продукции. Синтаксис задания продукций соответствует общепринятому, за исключением того, что вместо ε используется слово «epsilon» и любые символы грамматики должны быть разделены пробелами.

2.2 Класс Grammar

Класс `Grammar` содержит в себе всю информацию о грамматике, а также методы для работы с ней. Код конструктора класса представлен в листинге 1. Конструктор класса принимает на вход строку, содержащую описание грамматики и объект типа `Callable[[int, tuple[str, list[str]]], None]` – функцию-колбэк, которая будет вызываться при применении продукции. Эта функция должна принимать два аргумента: номер продукции и кортеж, содержащий левую и правую части продукции. Далее конструктор вызывает методы для парсинга продукций, нахождения терминалов, вычисления множеств FIRST и FOLLOW, заполнения таблицы синтаксического анализа.

Листинг 1. Код конструктора класса `Grammar`.

```
1  class Grammar:
2      EPSILON: str = "epsilon"
3
4      def __init__(
5          self,
6          text: str,
7          semantic_action: Callable[[int, tuple[str, list[str]]], None] | None = None,
8      ):
9          self productions: OrderedDict[str, list[list[str]]] = OrderedDict()
10         self.start_symbol: str = ""
11         self._parse_productions(text)
12
13         self.terminals: set[str] = set()
14         self._find_terminals()
15
16         self.first_sets: dict[str, set[str]] = {}
17         self._calculate_first_sets()
18
19         self.follow_sets: dict[str, set[str]] = {}
20         self._calculate_follow_sets()
21
22         self.lookup_table: dict[str, dict[str, list[str]]] = {}
23         self._fill_lookup_table()
24
25         # Сопоставляем уникальный номер с каждым правилом
26         self.rule_numbers = {}
27         rule_idx = 1
28         for nt, rules in self productions.items():
29             for rule in rules:
```

```

30         self.rule_numbers[(nt, tuple(rule))] = rule_idx
31         rule_idx += 1
32
33         # Semantic action callback
34         self.semantic_action = semantic_action

```

2.2.1 Метод `_parse productions`

Метод `_parse productions` выполняет разбор текстового описания грамматики и создаёт внутреннее представление продукций. Код метода представлен в листинге 2.

Метод принимает ссылку на объект класса `Grammar` и строку `text` типа `str`, содержащую текстовое представление грамматики.

Метод не возвращает значений явно, но заполняет словарь `productions` и устанавливает переменную `start_symbol` в объекте класса.

Листинг 2. Код метода `_parse productions`.

```

1  def _parse_productions(self, text: str):
2      for line in text.splitlines():
3          line = line.strip()
4          if not line:
5              continue
6
7          non_terminal, rule = line.split(">")
8          if self.start_symbol == "":
9              self.start_symbol = non_terminal.strip()
10
11         non_terminal = non_terminal.strip()
12         rules = [
13             [
14                 symbol.strip(' ')
15                 for symbol in re.findall(r"\.?\"|\S+", rule.strip())
16                 if symbol.strip(' ') != self.EPSILON
17             ]
18             for rule in rule.split("|")
19         ]
20
21         if non_terminal not in self.productions:
22             self.productions[non_terminal] = []
23
24         self.productions[non_terminal].extend(rules)

```

2.2.2 Метод `_calculate first sets`

Метод `_calculate first sets` вычисляет множества FIRST для всех символов грамматики. Код метода представлен в листинге 3.

Метод принимает ссылку на объект класса `Grammar`.

Метод не возвращает значений явно, но заполняет словарь `first_sets` в объекте класса.

Листинг 3. Код метода `_calculate_first_sets`.

```

1 def _calculate_first_sets(self):
2     # Инициализация FIRST для всех символов
3     for non_terminal in self productions:
4         self.first_sets[non_terminal] = set()
5
6     # Для терминалов FIRST содержит только сам терминал
7     for terminal in self.terminals:
8         self.first_sets[terminal] = {terminal}
9
10    # Вычисление FIRST для нетерминалов
11    changed = True
12    while changed:
13        changed = False
14        for non_terminal, rules in self productions.items():
15            for rule in rules:
16                if not rule: # Пустое правило эpsilon()
17                    if "" not in self.first_sets[non_terminal]:
18                        self.first_sets[non_terminal].add("")
19                        changed = True
20            else:
21                all_can_derive_epsilon = True
22                for i, symbol in enumerate(rule):
23                    # Добавляем все символы из FIRST(symbol) кроме эpsilon
24                    first_without_epsilon = self.first_sets[symbol] - {""}
25                    old_size = len(self.first_sets[non_terminal])
26                    self.first_sets[non_terminal].update(first_without_epsilon)
27                    if len(self.first_sets[non_terminal]) > old_size:
28                        changed = True
29
30                # Если symbol не может порождать эpsilon, прерываем
31                if "" not in self.first_sets[symbol]:
32                    all_can_derive_epsilon = False
33                    break
34
35                # Если все символы в правиле могут порождать эpsilon,
36                # то и нетерминал может порождать эpsilon
37                if (
38                    all_can_derive_epsilon
39                    and "" not in self.first_sets[non_terminal]
40                ):
41                    self.first_sets[non_terminal].add("")
42                    changed = True

```

2.2.3 Метод `_calculate_follow_sets`

Метод `_calculate_follow_sets` вычисляет множества FOLLOW для нетерминальных символов грамматики. Код метода представлен в листинге 4.

Метод принимает ссылку на объект класса Grammar.

Метод не возвращает значений явно, но заполняет словарь `follow_sets` в объекте класса.

Листинг 4. Код метода `_calculate_follow_sets`.

```

1 def _calculate_follow_sets(self):
2     # инициализировать  $Fo(S) = \{ \$ \}$ , а все остальные  $Fo(A_i)$  пустыми множествами
3     for non_terminal in self productions:
4         self.follow_sets[non_terminal] = set()
5
6     # Добавляем символ конца строки  $\$$  для начального символа
7     self.follow_sets[self.start_symbol].add("$")
8
9     # Повторяем, пока в наборах  $Follow$  происходят изменения
10    changed = True
11    while changed:
12        changed = False
13
14        # Для каждого нетерминала  $A_j$  в грамматике
15        for non_terminal_j, rules in self productions.items():
16
17            # Для каждого правила  $A_j \rightarrow w$ 
18            for rule in rules:
19
20                # Для каждого символа в правиле
21                for i, symbol in enumerate(rule):
22
23                    # Если символ — нетерминал  $A_i$ 
24                    if symbol in self productions:
25                        #  $w'$  — остаток правила после  $A_i$ 
26                        remainder = rule[i + 1 :] if i + 1 < len(rule) else []
27
28                        # Если есть терминалы после  $A_i$  ( $w'$  не пусто)
29                        if remainder:
30                            # Вычисляем  $First(w')$ 
31                            first_of_remainder = self._first_of_sequence(remainder)
32
33                            # Если терминал  $a$  находится в  $First(w')$ , то добавляем  $a$  к  $Follow(A_i)$ 
34                            for terminal in first_of_remainder - {" "}:
35                                if terminal not in self.follow_sets[symbol]:
36                                    self.follow_sets[symbol].add(terminal)
37                                    changed = True
38
39                            # Если  $\epsilon$  находится в  $First(w')$ , то добавляем  $Follow(A_j)$  к  $Follow(A_i)$ 
40                            if "" in first_of_remainder:
41                                old_size = len(self.follow_sets[symbol])
42                                self.follow_sets[symbol].update(
43                                    self.follow_sets[non_terminal_j]
44                                )
45                                if len(self.follow_sets[symbol]) > old_size:
46                                    changed = True
47
48                        # Если  $w'$  пусто ( $A_i$  в конце правила), то добавляем  $Follow(A_j)$  к  $Follow(A_i)$ 
49                        else:
50                            old_size = len(self.follow_sets[symbol])
51                            self.follow_sets[symbol].update(
52                                self.follow_sets[non_terminal_j]
53                            )
54                            if len(self.follow_sets[symbol]) > old_size:
55                                changed = True

```

2.2.4 Метод `_first_of_sequence`

Метод `_first_of_sequence` вычисляет множество FIRST для последовательности символов. Код метода представлен в листинге 5.

Метод принимает ссылку на объект класса `Grammar` и список символов `sequence`.

Метод возвращает множество (set) строк, представляющих FIRST для заданной последовательности.

Листинг 5. Код метода `_first_of_sequence`.

```
1 def _first_of_sequence(self, sequence):
2     """Вычисляет множество FIRST для последовательности символов"""
3     if not sequence:
4         return {""}
5
6     result = set()
7     all_can_derive_epsilon = True
8
9     for symbol in sequence:
10        # Добавляем First(symbol) без эpsilon к результату
11        result.update(self.first_sets[symbol] - {""})
12
13        # Если symbol не может породить эpsilon, останавливаемся
14        if "" not in self.first_sets[symbol]:
15            all_can_derive_epsilon = False
16            break
17
18    # Если все символы могут породить эpsilon, добавляем эpsilon к результату
19    if all_can_derive_epsilon:
20        result.add("")
21
22    return result
```

2.2.5 Метод `_fill_lookup_table`

Метод `_fill_lookup_table` заполняет таблицу синтаксического анализа для LL(1)-грамматики. Код метода представлен в листинге 6.

Метод принимает ссылку на объект класса `Grammar`.

Метод не возвращает значений явно, но заполняет словарь `lookup_table` в объекте класса или вызывает исключение, если грамматика не является LL(1).

Листинг 6. Код метода `_fill_lookup_table`.

```
1 def _fill_lookup_table(self):
2     # Для каждого нетерминала A
3     for non_terminal, rules in self productions.items():
4         # Формируем таблицу синтаксического анализа
5         self.lookup_table[non_terminal] = {}
6
7         # Для каждого правила A → w
8         for rule in rules:
9             # Вычисляем First(w)
```

```

10     first_of_rule = self._first_of_sequence(rule)
11
12     # Для каждого терминала  $a$  в  $First(w)$ 
13     for terminal in first_of_rule - {" "}:
14         # Если  $T[A, a]$  уже содержит правило, грамматика не  $LL(1)$ 
15         if terminal in self.lookup_table[non_terminal]:
16             raise ValueError(
17                 "Грамматика не является  $LL(1)$  грамматикой—.\n"
18                 f"Конфликт в ячейке_{non_terminal}_{terminal}.\n"
19                 f"Новое правило:_{non_terminal}_{terminal} → {self.lookup_table[non_terminal][terminal]}\n"
20                 f"Существующее правило:_{non_terminal}_{terminal} → {self.lookup_table[non_terminal][terminal]}\n"
21             )
22
23     # Добавляем правило в таблицу
24     self.lookup_table[non_terminal][terminal] = rule
25
26     # Если эpsilon в  $First(w)$ , то для каждого  $b$  в  $Follow(A)$ 
27     if "" in first_of_rule:
28         for terminal in self.follow_sets[non_terminal]:
29             # Если  $T[A, b]$  уже содержит правило, грамматика не  $LL(1)$ 
30             if terminal in self.lookup_table[non_terminal]:
31                 raise ValueError(
32                     "Грамматика не является  $LL(1)$  грамматикой—.\n"
33                     f"Конфликт в ячейке_{non_terminal}_{terminal}.\n"
34                     f"Новое правило:_{non_terminal}_{terminal} → {self.lookup_table[non_terminal][terminal]}\n"
35                     f"Существующее правило:_{non_terminal}_{terminal} → {self.lookup_table[non_terminal][terminal]}\n"
36                 )
37
38     # Добавляем правило в таблицу
39     self.lookup_table[non_terminal][terminal] = rule

```

2.2.6 Метод format_rules

Метод `format_rules` форматирует все правила грамматики для удобного представления. Код метода представлен в листинге 7.

Метод принимает ссылку на объект класса `Grammar`.

Метод возвращает строку (`str`), содержащую все правила грамматики с их номерами.

Листинг 7. Код метода `format_rules`.

```

1 def format_rules(self) -> str:
2     result = []
3     sorted_rules = sorted(self.rule_numbers.items(), key=lambda x: x[1])
4
5     for rule, number in sorted_rules:
6         non_terminal, symbols = rule
7         rule_text = f"{number}:_{non_terminal}_{symbols} → {self.lookup_table[non_terminal][symbols]}\n"
8         result.append(rule_text)
9
10    return "\n".join(result)

```

2.2.7 Метод `format_lookup_table`

Метод `format_lookup_table` форматирует таблицу синтаксического анализа для удобного представления. Код метода представлен в листинге 8.

Метод принимает ссылку на объект класса `Grammar`.

Метод возвращает строку (`str`), содержащую таблицу синтаксического анализа в виде таблицы с использованием библиотеки `PrettyTable`.

Листинг 8. Код метода `format_lookup_table`.

```
1 def format_lookup_table(self) -> str:
2     table = PrettyTable()
3     terminals = list(self.terminals)
4     table.field_names = [""] + terminals + ["$"]
5
6     for non_terminal in self productions:
7         row = [non_terminal]
8         for terminal in terminals + ["$"]:
9             if terminal in self.lookup_table[non_terminal]:
10                 rule = self.lookup_table[non_terminal][terminal]
11                 rule_num = self.rule_numbers.get((non_terminal, tuple(rule)), "")
12                 row.append(f"{rule_num}:_{'_'}.join(rule)}")
13             else:
14                 row.append("_-_-")
15         table.add_row(row)
16     return str(table)
```

2.2.8 Метод `format_first_sets`

Метод `format_first_sets` форматирует множества FIRST для удобного представления. Код метода представлен в листинге 9.

Метод принимает ссылку на объект класса `Grammar`.

Метод возвращает строку (`str`), содержащую множества FIRST для всех символов грамматики.

Листинг 9. Код метода `format_first_sets`.

```
1 def format_first_sets(self) -> str:
2     """Форматирует множества FIRST в читаемый вид. """
3     result = []
4     result.append("Множества_FIRST:")
5     result.append("=" * 40)
6
7     # Сортируем для гарантии порядка вывода
8     for symbol in sorted(self.first_sets.keys()):
9         # Заменяем пустую строку на эпсилон для лучшей читаемости
10        first_set = {
11            self.EPSILON if item == "" else item for item in self.first_sets[symbol]
12        }
13        result.append(f"FIRST({symbol})_{'_'}.join(sorted(first_set))}")
14
15    return "\n".join(result)
```

2.2.9 Метод `format_follow_sets`

Метод `format_follow_sets` форматирует множества FOLLOW для удобного представления. Код метода представлен в листинге 10.

Метод принимает ссылку на объект класса `Grammar`.

Метод возвращает строку (`str`), содержащую множества FOLLOW для всех нетерминалов грамматики.

Листинг 10. Код метода `format_follow_sets`.

```
1 def format_follow_sets(self) -> str:
2     """Форматирует множества FOLLOW в читаемый вид."""
3     result = []
4     result.append("Множества FOLLOW:")
5     result.append("=" * 40)
6
7     # Обработываем только нетерминалы
8     for non_terminal in sorted(self productions.keys()):
9         follow_set = self.follow_sets.get(non_terminal, set())
10        result.append(
11            f"FOLLOW({non_terminal}) = {', '.join(sorted(follow_set))}"
12        )
13
14    return "\n".join(result)
```

2.2.10 Метод `analyze`

Метод `analyze` выполняет синтаксический анализ входной последовательности токенов с использованием LL(1)-алгоритма. После применения продукции метод вызывает функцию-колбэк, представляющую семантические действия, если она была задана при создании объекта класса `Grammar`. Код метода представлен в листинге 11.

Метод принимает ссылку на объект класса `Grammar` и список `input_tokens` типа `list[str]`, представляющий последовательность входных токенов.

Метод возвращает список (`list`) целых чисел, содержащий номера применённых правил в процессе анализа.

Листинг 11. Код метода `analyze`.

```
1 def analyze(self, input_tokens: list[str]) -> list[int]:
2     input_tokens = input_tokens.copy()
3     input_tokens += ["$"]
4     input_pos = 0
5
6     # Инициализируем стек с терминальным символом и начальным символом
7     stack = ["$", self.start_symbol]
8
9     rules_applied = []
10
11    while stack:
12        top = stack[-1]
13
```

```

14     current_symbol = (
15         input_tokens[input_pos] if input_pos < len(input_tokens) else "$"
16     )
17
18     # Случай 1: Верхний символ стека — нетерминал
19     if top in self productions:
20         # Ищем правило в таблице синтаксического анализа
21         if current_symbol in self.lookup_table[top]:
22             # Получаем правило
23             production = self.lookup_table[top][current_symbol]
24
25             # Удаляем нетерминал из стека
26             stack.pop()
27
28             # Добавляем правило в rules_applied
29             rule_number = self.rule_numbers[(top, tuple(production))]
30             rules_applied.append(rule_number)
31
32             # Execute semantic action if provided
33             if self.semantic_action:
34                 self.semantic_action(rule_number, (top, production))
35
36             # Добавляем правило в стек в обратном порядке
37             for symbol in reversed(production):
38                 stack.append(symbol)
39         else:
40             expected_symbols = list(self.lookup_table[top].keys())
41             raise ValueError(
42                 f"Syntax_error: _expected_one_of_{expected_symbols}, _got_{current_symbol}"
43             )
44
45     # Случай 2: Верхний символ стека — терминал
46     elif top != "$":
47         if top == current_symbol:
48             # Удаляем терминал из стека
49             stack.pop()
50             # Переходим к следующему символу ввода
51             input_pos += 1
52         else:
53             raise ValueError(
54                 f"Syntax_error: _expected_{top}, _got_{current_symbol}"
55             )
56
57     # Случай 3: Верхний символ стека — $
58     else: # top == "$"
59         if current_symbol == "$":
60             # Успешный синтаксический анализ
61             stack.pop() # Удаляем $ из стека
62         else:
63             raise ValueError(
64                 f"Syntax_error: _unexpected_symbols_at_end_of_input:_{current_symbol}"
65             )
66
67     return rules_applied

```

2.2.11 Метод generate

Метод **generate** генерирует случайное предложение по заданной грамматике. Код метода представлен в листинге 12.

Метод принимает ссылку на объект класса **Grammar** и необязательный параметр **symbol** типа **str** или **None**, указывающий начальный символ для генерации.

Метод возвращает кортеж (**tuple**), содержащий список терминалов и список номеров применённых правил.

Листинг 12. Код метода **generate**.

```
1 def generate(self, symbol: str | None = None) -> tuple[list[str], list[int]]:
2     """Генерирует предложение по заданной грамматике. Возвращает список терминалов
3     и список номеров применённых правил. """
4
5     if symbol is None:
6         return self.generate(self.start_symbol)
7
8     # Если символ — терминал, возвращаем его
9     if symbol not in self productions:
10         return [symbol], []
11
12     # Выбираем случайное правило для нетерминала
13     rules = self productions[symbol]
14     chosen_rule = random.choice(rules)
15
16     # Получаем номер выбранного правила
17     rule_number = self.rule_numbers[(symbol, tuple(chosen_rule))]
18
19     # Инициализируем результаты
20     terminals = []
21     rule_numbers = [rule_number]
22
23     # Разворачиваем каждый символ в правой части правила
24     for s in chosen_rule:
25         sub_terminals, sub_rules = self.generate(s)
26         terminals.extend(sub_terminals)
27         rule_numbers.extend(sub_rules)
28
29     return terminals, rule_numbers
```

2.2.12 Метод generate_derivation_steps

Метод **generate_derivation_steps** преобразует список номеров правил в последовательность шагов вывода. Код метода представлен в листинге 13.

Метод принимает ссылку на объект класса **Grammar** и список **rule_numbers** типа **list[int]**, содержащий номера правил для применения.

Метод возвращает список (**list**) строк, представляющий каждый шаг вывода предложения.

Листинг 13. Код метода **generate_derivation_steps**.

```

1 def generate_derivation_steps(self, rule_numbers: list[int]) -> list[str]:
2     """Преобразует список номеров правил в последовательность шагов вывода.
3     Возвращает список строк, представляющих каждый шаг вывода. """
4
5     # Получаем соответствие между номерами правил и самими правилами
6     rule_details = {num: rule for rule, num in self.rule_numbers.items()}
7
8     # Начинаем с начального символа
9     current = self.start_symbol
10    steps = [current]
11
12    # Применяем каждое правило по порядку
13    for rule_num in rule_numbers:
14        if rule_num in rule_details:
15            non_terminal, replacement = rule_details[rule_num]
16
17            # Находим первое вхождение нетерминала и заменяем его
18            words = current.split()
19            for i, word in enumerate(words):
20                if word == non_terminal:
21                    words[i : i + 1] = replacement
22                    break
23
24            current = "_".join(words)
25            steps.append(current)
26
27    return steps

```

2.3 Функция load_grammar

Функция `load_grammar` загружает грамматику из файла и возвращает объект класса `Grammar`. Функция принимает необязательный параметр `filename` типа `str`, указывающий имя файла с грамматикой. Функция возвращает `None` в случае ошибки и объект класса `Grammar` в случае успешной загрузки грамматики. Код функции представлен в листинге 14. Также в ней задаются семантические действия, которые будут выполняться при применении продукций, и сохраняются в файлы множества `FIRST` и `FOLLOW` и таблица синтаксического анализа.

Листинг 14. Код функции `load_grammar`.

```

1 def load_grammar(filename: str = "grammar.txt") -> Grammar | None:
2     try:
3         # #b - 2 * #a + 3 * #c
4         actions = [
5             lambda rule_number, applied_count, _: print(
6                 f"Rule_{rule_number}_{applied_count}_{applied_count}_times):_iconst_1"
7             ),
8             lambda rule_number, applied_count, _: print(
9                 f"Rule_{rule_number}_{applied_count}_{applied_count}_times):_iconst_2_sub"
10            ),
11            lambda rule_number, applied_count, _: print(
12                f"Rule_{rule_number}_{applied_count}_{applied_count}_times):_iconst_1_iadd"
13            ),
14            lambda rule_number, applied_count, _: print(

```

```

15         f"Rule_{rule_number}_{applied_x{applied_count}_times}:_const_1_iadd"
16     ),
17     lambda rule_number, applied_count, _: print(
18         f"Rule_{rule_number}_{applied_x{applied_count}_times}:_const_3_iadd"
19     ),
20 ]
21
22 with open(filename, "r", encoding="utf-8") as file:
23     text = file.read()
24 grammar = Grammar(text, ActionsListWithAppliedCount(actions))
25
26 # Сохраняем информацию о грамматике в файлы
27 with open("grammar_rules.txt", "w", encoding="utf-8") as output_file:
28     output_file.write(grammar.format_rules())
29     print("Правила грамматики с номерами сохранены в grammar_rules.txt")
30
31 with open("grammar_lookup_table.txt", "w", encoding="utf-8") as output_file:
32     output_file.write(grammar.format_lookup_table())
33     print(
34         "Таблица синтаксического анализа сохранена в grammar_lookup_table.txt"
35     )
36
37 with open("grammar_first.txt", "w", encoding="utf-8") as output_file:
38     output_file.write(grammar.format_first_sets())
39     print("Множества FIRST сохранены в grammar_first.txt")
40
41 with open("grammar_follow.txt", "w", encoding="utf-8") as output_file:
42     output_file.write(grammar.format_follow_sets())
43     print("Множества FOLLOW сохранены в grammar_follow.txt")
44
45     print(f"Грамматика успешно загружена из файла {filename}")
46     return grammar
47 except FileNotFoundError:
48     print(f"Ошибка: Файл {filename} не найден")
49     return None
50 except ValueError as e:
51     print(f"Ошибка при загрузке грамматики: {e}")
52     return None
53 except Exception as e:
54     print(f"Неизвестная ошибка: {e}")
55     return None

```

2.4 Функция main

Функция `main` реализует интерактивную консольную оболочку для взаимодействия с пользователем. Код функции представлен в листинге 15.

Листинг 15. Код функции `main`.

```

1 def main():
2     print("Программа для работы с LL(1) грамматиками")
3     print("=" * 60)
4     print("Варианты команд:")
5     print("_ _ load_файл <> _ _ загрузить грамматику из файла по умолчанию grammar.txt")

```

```

6  print("_check_строка<>_проверить,_соответствует_ли_строка_грамматике")
7  print("_generate_сгенерировать_случайную_строку_по_грамматике")
8  print("_exit_выход_из_программы")
9  print("=" * 60)
10
11  # Загружаем грамматику по умолчанию при старте
12  grammar = load_grammar()
13
14  while True:
15      command = input("\Введите команду:_").strip()
16
17      if not command:
18          continue
19
20      parts = command.split(maxsplit=1)
21      cmd = parts[0].lower()
22
23      if cmd == "exit":
24          print("Выход_из_программы.")
25          break
26
27      elif cmd == "load":
28          filename = "grammar.txt"
29          if len(parts) > 1:
30              filename = parts[1].strip()
31          grammar = load_grammar(filename)
32
33      elif cmd == "check":
34          input_string = ""
35          if len(parts) > 1:
36              input_string = parts[1].strip()
37          check_string(grammar, input_string)
38
39      elif cmd == "generate":
40          generate_string(grammar)
41
42      else:
43          print(f"Неизвестная команда:_{cmd}")
44          print("Доступные команды:_load,_check,_generate,_exit")

```

3 Результаты работы программы

Результаты работы программы представлены на Рис. 3.

```
Программа для работы с LL(1)-грамматиками
=====
Варианты команд:
- load <файл> - загрузить грамматику из файла (по умолчанию grammar.txt)
- check <строка> - проверить, соответствует ли строка грамматике
- generate - сгенерировать случайную строку по грамматике
- exit - выход из программы
=====
Правила грамматики с номерами сохранены в grammar_rules.txt
Таблица синтаксического анализа сохранена в grammar_lookup_table.txt
Множества FIRST сохранены в grammar_first.txt
Множества FOLLOW сохранены в grammar_follow.txt
Грамматика успешно загружена из файла grammar.txt

Введите команду: generate
Сгенерированная строка: b a b b b c b b b
Применённые правила: [1, 4, 2, 3, 4, 4, 5, 3, 4]
Подробный результат генерации сохранен в generation_result.txt

Введите команду: check b b b b
Проверка строки: 'b b b b'
Rule #1 (applied x1 times): iconst_1
Rule #4 (applied x1 times): iconst_1 iadd
Rule #3 (applied x1 times): iconst_1 iadd
Rule #4 (applied x2 times): iconst_1 iadd
Результат: Строка соответствует грамматике
Применённые правила: [1, 4, 3, 4]
Подробный результат анализа сохранен в analysis_result.txt

Введите команду: check
Проверка строки: ''
Результат: Строка не соответствует грамматике
Ошибка: Syntax error: expected one of ['b'], got '$'

Введите команду: exit
Выход из программы.
```

Рис. 3. Результаты работы программы.

```

Введите команду: abrakadabra
Неизвестная команда: abrakadabra
Доступные команды: load, check, generate, exit

Введите команду: load nonexistent.txt
Ошибка: Файл nonexistent.txt не найден

```

Рис. 4. Реакция программы на некорректный пользовательский ввод.

На Рис. 4 представлена реакция программы на некорректный пользовательский ввод.

Листинг 16. Содержимое файла `grammar_first.txt`.

```

1 Множества
2 FIRST:
3 =====
4 FIRST(A) = {a, b}
5 FIRST(B) = {b}
6 FIRST(C) = {c}
7 FIRST(S) = {b}
8 FIRST(a) = {a}
9 FIRST(b) = {b}
10 FIRST(c) = {c}

```

Листинг 17. Содержимое файла `grammar_follow.txt`.

```

1 Множества
2 FOLLOW:
3 =====
4 FOLLOW(A) = {b}
5 FOLLOW(B) = {a, b, c}
6 FOLLOW(C) = {b}
7 FOLLOW(S) = {$}

```

Листинг 18. Содержимое файла `grammar_lookup_table.txt`.

```

1 +---+-----+-----+-----+-----+
2 |   |      a      |      b      |      c      |      $      |
3 +---+-----+-----+-----+-----+
4 | S |      -      | 1: B A b |      -      |      -      |
5 | A | 2: a A B C | 3: b B   |      -      |      -      |
6 | B |      -      | 4: b    |      -      |      -      |
7 | C |      -      |      -   | 5: c A   |      -      |
8 +---+-----+-----+-----+-----+

```

В листингах 16, 17 и 18 представлено содержимое файлов `grammar_first.txt`, `grammar_follow.txt` и `grammar_lookup_table.txt` соответственно.

Заключение

В ходе выполнения лабораторной работы была реализована программа – синтаксический анализатор для LL(1)-грамматик. С его помощью для заданной грамматики были построены множества FIRST и FOLLOW для всех нетерминалов и таблица синтаксического анализа. Программа-анализатор также включает в себя функции для проверки принадлежности цепочки языку, заданному грамматикой, и генерации случайных цепочек. Также программа позволяет задать семантические действия, которые будут выполняться при применении продукций. Для демонстрации возможностей программы в качестве семантических действий решено было использовать генерацию инструкций для Java Virtual Machine.

Из достоинств выполнения лабораторной работы можно выделить, во-первых, возможность задания грамматики в отдельном текстовом файле, что позволяет легко изменять и расширять её без модификации программного кода. Во-вторых, программа автоматически проверяет, что введенная грамматика является LL(1)-грамматикой, и вычисляет множества FIRST и FOLLOW для всех нетерминалов, а также таблицу синтаксического анализа. В противном случае, программа выводит сообщение об ошибке, в котором указывается на конкретные правила грамматики, между выбором которых возникает неоднозначность. В третьих, программа позволяет легко задавать произвольные семантические действия в виде функций или объектов колбэков, которые будут вызываться при применении продукций.

К недостаткам текущей реализации можно отнести то, что алгоритм генерации не контролирует длину предложений, что может приводить к избыточно длинным или коротким цепочкам.

Функционал программы несложно масштабировать. Например, несложно добавить функцию для визуализации дерева вывода распознаваемой цепочки, пусть даже в текстовом виде, в классе Grammar уже есть все необходимые для этого данные. Кроме того, класс Grammar может служить хорошей основой для создания полноценного LL(k) анализатора.

На выполнение лабораторной работы ушло около 4 часов, основная часть кода была взята из лабораторной работы №3. Работа была выполнена в среде разработки Visual Studio Code. Программа написана на Python версии 3.13.

Список литературы

- [1] Востров, А.В. Курс лекций по дисциплине «Математическая логика». URL <https://tema.spbstu.ru/compiler/> (дата обращения 01.04.2025 г.)
- [2] Лутц, М. Изучаем Python. 5-е изд. / М. Лутц. — СПб.: Питер, 2019. — 1216 с.