

# Приложения теории формальных языков и синтаксического анализа

Семён Григорьев

30 июля 2021 г.



# Оглавление

|          |                                                                           |           |
|----------|---------------------------------------------------------------------------|-----------|
| <b>1</b> | <b>Некоторые понятия линейной алгебры</b>                                 | <b>13</b> |
| 1.1      | Бинарные операции и их свойства . . . . .                                 | 13        |
| 1.2      | Полугруппа . . . . .                                                      | 16        |
| 1.3      | Моноид . . . . .                                                          | 17        |
| 1.4      | Группа . . . . .                                                          | 17        |
| 1.5      | Полукольцо . . . . .                                                      | 18        |
| 1.6      | Кольцо . . . . .                                                          | 20        |
| 1.7      | Матрицы и вектора . . . . .                                               | 20        |
| 1.8      | Теоретическая сложность умножения матриц . . . . .                        | 24        |
| <b>2</b> | <b>Некоторые сведения из теории графов</b>                                | <b>29</b> |
| 2.1      | Основные определения . . . . .                                            | 29        |
| 2.2      | Задачи поиска путей . . . . .                                             | 34        |
| 2.3      | Алгоритм Флойда-Уоршелла . . . . .                                        | 36        |
| 2.4      | Анализ путей в графе и линейная алгебра . . . . .                         | 40        |
| <b>3</b> | <b>Общие сведения теории формальных языков</b>                            | <b>45</b> |
| <b>4</b> | <b>Задача о поиске путей с ограничениями в терминах формальных языков</b> | <b>47</b> |
| 4.1      | Постановка задачи . . . . .                                               | 47        |
| 4.2      | О разрешимости задачи . . . . .                                           | 49        |
| 4.3      | Области применения . . . . .                                              | 50        |

|           |                                                                  |            |
|-----------|------------------------------------------------------------------|------------|
| <b>5</b>  | <b>Контекстно-свободные грамматики и языки</b>                   | <b>51</b>  |
| 5.1       | Дерево вывода . . . . .                                          | 54         |
| 5.2       | Пустота КС-языка . . . . .                                       | 55         |
| 5.3       | Нормальная форма Хомского . . . . .                              | 56         |
| 5.4       | Лемма о накачке . . . . .                                        | 60         |
| 5.5       | Замкнутость КС языков относительно операций . . . . .            | 62         |
| <b>6</b>  | <b>СҮК для вычисления КС запросов</b>                            | <b>65</b>  |
| 6.1       | Алгоритм СҮК . . . . .                                           | 65         |
| 6.2       | Алгоритм для графов на основе СҮК . . . . .                      | 69         |
| <b>7</b>  | <b>КС и конъюнктивная достижимость через произведение матриц</b> | <b>81</b>  |
| 7.1       | КС достижимость через произведение матриц . . . . .              | 81         |
| 7.1.1     | Расширение алгоритма для конъюнктивных грамматик . . . . .       | 85         |
| 7.2       | Особенности реализации . . . . .                                 | 87         |
| <b>8</b>  | <b>КС достижимость через тензорное произведение</b>              | <b>91</b>  |
| 8.1       | Рекурсивные автоматы и сети . . . . .                            | 91         |
| 8.2       | Тензорное произведение . . . . .                                 | 92         |
| 8.3       | Алгоритм . . . . .                                               | 94         |
| 8.4       | Примеры . . . . .                                                | 97         |
| 8.5       | Особенности реализации . . . . .                                 | 115        |
| <b>9</b>  | <b>Сжатое представление леса разбора</b>                         | <b>117</b> |
| 9.1       | Лес разбора как представление контекстно-свободной грамматики    | 117        |
| <b>10</b> | <b>Алгоритм на основе нисходящего анализа</b>                    | <b>129</b> |
| 10.1      | Рекурсивный спуск . . . . .                                      | 129        |
| 10.2      | LL(k)-алгоритм синтаксического анализа . . . . .                 | 130        |
| 10.3      | Алгоритм Generalized LL . . . . .                                | 136        |
| 10.4      | Алгоритм вычисления КС запросов на основе GLL . . . . .          | 144        |

|                                                     |            |
|-----------------------------------------------------|------------|
| <b>11 Алгоритм на основе восходящего анализа</b>    | <b>147</b> |
| 11.1 Восходящий синтаксический анализ . . . . .     | 147        |
| 11.1.1 LR(0) алгоритм . . . . .                     | 149        |
| 11.1.2 SLR(1) алгоритм . . . . .                    | 150        |
| 11.1.3 CLR(1) алгоритм . . . . .                    | 151        |
| 11.1.4 Примеры . . . . .                            | 151        |
| 11.1.5 Сравнение классов LL и LR . . . . .          | 158        |
| 11.2 GLR и его применение для КС запросов . . . . . | 158        |
| 11.2.1 Классический GLR алгоритм . . . . .          | 159        |
| 11.2.2 Модификации GLR . . . . .                    | 163        |



# Список авторов

- **Семён Григорьев**

Санкт-Петербургский государственный университет, Университетская набережная, 7/9, Санкт-Петербург, 199034, Россия

`s.v.grigoriev@spbu.ru`

JetBrains Research, Приморский проспект 68-70, здание 1, Санкт-Петербург, 197374, Россия

`semyon.grigorev@jetbrains.com`

- **Екатерина Вербицкая**

JetBrains Research, Приморский проспект 68-70, здание 1, Санкт-Петербург, 197374, Россия

`ekaterina.verbitskaya@jetbrains.com`

- **Дмитрий Кутленков**

Санкт-Петербургский государственный университет, Университетская набережная, 7/9, Санкт-Петербург, 199034, Россия

`kutlenkov.dmitri@gmail.com`





# Введение

Теория формальных языков находит применение не только для ставших уже классическими задач синтаксического анализа кода (языков программирования, искусственных языков) и естественных языков, но и в других областях, таких как статический анализ кода, графовые базы данных, биоинформатика, машинное обучение.

Например, в машинном обучении использование формальных грамматик позволяет передать искусственной генеративной нейронной сети, предназначенной для построения цепочек с определёнными свойствами, знания о синтаксической структуре этих цепочек, что позволяет существенно упростить процесс обучения и повысить качество результата [41]. Вместе с этим, развиваются подходы, позволяющие нейронным сетям наоборот извлекать синтаксическую структуру (строить дерево вывода) для входных цепочек [37, 38].

В биоинформатике формальные грамматики нашли широкое применение для описания особенностей вторичной структуры геномных и белковых последовательностей [25, 4, 80]. Соответствующие алгоритмы синтаксического анализа используются при создании инструментов обработки данных.

Таким образом, теория формальных языков выступает в качестве основы для многих прикладных областей, а алгоритмы синтаксического анализа применимы не только для обработки естественных языков или языков программирования. Нас же в данной работе будет интересовать применение теории формальных языков и алгоритмов синтаксического анализа для анализа графовых баз данных и для статического анализа кода.

Одна из классических задач, связанных с анализом графов — это поиск путей в графе. Возможны различные формулировки этой задачи. В некоторых случаях необходимо выяснить, существует ли путь с определёнными свойствами между двумя выбранными вершинами. В других же ситуациях необходимо найти все пути в графе, удовлетворяющие некоторым свойствам или ограничениям. Например, в качестве ограничений можно указать, что искомым путь должен быть простым, кратчайшим, гамильтоновым и так далее.

Один из способов задавать ограничения на пути в графе основан на использовании формальных языков. Базовое определение языка говорит нам, что язык — это множество слов над некоторым алфавитом. Если рассмотреть граф, рёбра которого помечены символами из алфавита, то путь в таком графе будет задавать слово: достаточно соединить последовательно символы, лежащие на рёбрах пути. Множество же таких путей будет задавать множество слов или язык. Таким образом, если мы хотим найти некоторое множество путей в графе, то в качестве ограничения можно описать язык, который должно задавать это множество. Иными словами, задача поиска путей может быть сформулирована следующим образом: необходимо найти такие пути в графе, что слова, получаемые конкатенацией меток их рёбер, принадлежат заданному языку. Такой класс задач будем называть задачами поиска путей с ограничениям в терминах формальных языков.

Подобный класс задач часто возникает в областях, связанных с анализом граф-структурированных данных и активно исследуется [11, 5, 34, 70, 10, 11]. Исследуются как классы языков, применяемых для задания ограничений, так и различные постановки задачи.

Граф-структурированные данные встречаются не только в графовых базах данных, но и при статическом анализе кода: по программе можно построить различные графы отображающие её свойства. Скажем, граф вызовов, граф потока данных и так далее. Оказывается, что поиск путей в специального вида графах с использованием ограничений в терминах формальных языков позволяет исследовать некоторые нетривиальные свойства программы. Например проводить межпроцедурный анализ указателей или анализ алиасов [79, 74, 77], строить срезы программ [52], проводить анализ типов [49].

В данной работе представлен ряд алгоритмов для поиска путей с ограничениями в терминах формальных языков. Основной акцент будет сделан на контекстно-свободных языках, однако будут затронуты и другие классы: регулярные, многокомпонентные контекстно-свободные (Multiple Context-Free Languages, MCFL [59]) и конъюнктивные языки. Будет показано, что теория формальных языков и алгоритмы синтаксического анализа применимы не только для анализа языков программирования или естественных языков, а также для анализа графовых баз данных и статического анализа кода, что приводит к возникновению новых задач и переосмыслению старых.

Структура данной работы такова. В начале (в части 2) мы рассмотрим основные понятия из теории графов, необходимые в данной работе. Данные разделы являются подготовительными и не обязательны к прочтению, если такие понятия как *ориентированный граф* и *матрица смежности* уже известны читателю. Более того, они лишь вводят определения, подразумевая,

что более детальное изучение соответствующих разделов науки остается за рамками этой работы и скорее всего уже проделано читателем. Затем, в главе 3 мы введём основные понятия из теории формальных языков. Далее, в главе 4 рассмотрим различные варианты постановки задачи поиска путей с ограничениями в терминах формальных языков, обсудим базовые свойства задач, её разрешимость в различных постановках и т.д.. И в итоге зафиксируем постановку, которую будем изучать далее. После этого, в главах 6–11 мы будем подробно рассматривать различные алгоритмы решения этой задачи, попутно вводя специфичные для рассматриваемого алгоритма структуры данных. Большинство алгоритмов будут основаны на классических алгоритмах синтаксического анализа, таких как СΥК или LR.



# Глава 1

## Некоторые понятия линейной алгебры<sup>1</sup>

При изложении ряда алгоритмов будут активно использоваться некоторые понятия и инструменты линейной алгебры, такие как моноид, полукольцо или матрица. В данном разделе необходимые понятия будут определены и приведены некоторые примеры соответствующих конструкций. Для более глубокого изучения материала рекомендуются обратиться к соответствующим разделам алгебры.

### 1.1 Бинарные операции и их свойства

Введём понятие *бинарной операции* и рассмотрим некоторые её свойства, такие как *коммутативность* и *ассоциативность*.

**Определение 1.1.1.** Функцию, принимающую два аргумента,  $f : S \times K \rightarrow Q$  будем называть *двухместной* или *функцией arity два*. Для записи таких функций будем использовать типичную нотацию:  $c = f(a, b)$ .

**Определение 1.1.2.** *Бинарная операция* — это двухместная функция, от которой дополнительно требуется, чтобы оба аргумента и результат лежали в

---

<sup>1</sup>Необходимо понимать, что, с одной стороны, в данном разделе рассматриваются самые базовые понятия, которые даются практически в любом учебнике алгебры. С другой же стороны, определения данных понятий оказываются весьма вариативными и часто вызывают дискуссии. Например, интересный анализ тонкостей определения группы можно найти в первом и втором параграфах первого раздела книги Николая Александровича Вавилова “Конкретная теория групп” [83]. Мы же дадим определения, удобные для дальнейшего изложения материала.

одном и том же множестве:  $f : S \times S \rightarrow S$ . В таком случае говорят, что бинарная операция определена на некотором множестве  $S$ . Для обозначения произвольной бинарной операции будем использовать символ  $\circ$  и пользоваться инфиксной нотацией для записи:  $c = a \circ b$ .

**Определение 1.1.3.** *Внешняя бинарная операция* — это бинарная операция, у которой аргументы лежат в разных множествах, при этом результат — в одном из этих множеств. Иными словами  $\circ : K \times S \rightarrow S$ , где  $K$  может не совпадать с  $S$  — внешняя бинарная операция.

Необходимо помнить, что как функции, так и бинарные операции, могут быть частично определёнными (частичные функции, частичные бинарные операции). Типичным примером частично определённой бинарной операции является деление на целых числах: она не определена, если второй аргумент равен нулю.

Бинарные операции могут обладать некоторыми дополнительными свойствами, такими как *коммутативность* или *ассоциативность*, позволяющими преобразовывать выражения, составленные с использованием этих операций.

**Определение 1.1.4.** Бинарная операция  $\circ : S \times S \rightarrow S$  называется *коммутативной*, если для любых  $x_1 \in S, x_2 \in S$  верно, что  $x_1 \circ x_2 = x_2 \circ x_1$ .

**Пример 1.1.1.** Рассмотрим несколько примеров коммутативных и некоммутативных операций.

- Операция сложения на целых числах  $+$  является коммутативной: известный ещё со школы перестановочный закон сложения.
- Операция конкатенации на строках  $\cdot$  не является коммутативной:

$$“ab” \cdot “c” = “abc” \neq “c” \cdot “ab” = “cab”.$$

- Операция умножения на целых числах является коммутативной: известный ещё со школы перестановочный закон умножения.
- Операция умножения матриц (над целыми числами)  $\cdot$  не является коммутативной:

$$\begin{pmatrix} 1 & 1 \\ 0 & 0 \end{pmatrix} \cdot \begin{pmatrix} 0 & 0 \\ 1 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 1 \\ 0 & 0 \end{pmatrix} \neq \begin{pmatrix} 0 & 0 \\ 1 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 1 \\ 0 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 0 \\ 1 & 1 \end{pmatrix}.$$

**Определение 1.1.5.** Бинарная операция  $\circ : S \times S \rightarrow S$  называется *ассоциативной*, если для любых  $x_1 \in S, x_2 \in S, x_3 \in S$  верно, что  $(x_1 \circ x_2) \circ x_3 = x_1 \circ (x_2 \circ x_3)$ . Иными словами, для ассоциативной операции результат вычислений не зависит от порядка применения операций.

**Пример 1.1.2.** Рассмотрим несколько примеров ассоциативных и неассоциативных операций.

- Операция сложения на целых числах  $+$  является ассоциативной.
- Операция умножения на целых числах является ассоциативной.
- Операция конкатенации на строках  $\cdot$  является ассоциативной:

$$("a" \cdot "b") \cdot "c" = "a" \cdot ("b" \cdot "c") = "abc".$$

- Операция возведения в степень (над целыми числами)  $^{\wedge}$  не является ассоциативной:

$$(2^{\wedge}2)^{\wedge}3 = 4^{\wedge}3 = 64 \neq 2^{\wedge}(2^{\wedge}3) = 2^{\wedge}8 = 256.$$

**Определение 1.1.6.** Говорят, что бинарная операция  $\otimes : S \times S \rightarrow S$  является *дистрибутивной* относительно бинарной операции  $\oplus : S \times S \rightarrow S$ , если

1. Для любых  $x_1, x_2, x_3 \in S, x_1 \otimes (x_2 \oplus x_3) = (x_1 \otimes x_2) \oplus (x_1 \otimes x_3)$  (дистрибутивность слева).
2. Для любых  $x_1, x_2, x_3 \in S, (x_2 \oplus x_3) \otimes x_1 = (x_2 \otimes x_1) \oplus (x_3 \otimes x_1)$  (дистрибутивность справа).

Если операция  $\otimes$  является коммутативной, то дистрибутивность слева и справа равносильны.

**Пример 1.1.3.** Рассмотрим несколько примеров дистрибутивных операций.

- Умножение целых чисел дистрибутивно относительно сложения и вычитания: классический *распределительный закон*, знакомый всем со школы.
- Операция деления (допустим, на действительных числах) не коммутативна. При этом, она дистрибутивна справа относительно сложения и вычитания, но не дистрибутивна слева.

$$(a + b)/c = (a/c) + (b/c)$$

но

$$c/(a+b) \neq (c/a) + (c/b)^2.$$

**Определение 1.1.7.** Бинарная операция  $\circ : S \times S \rightarrow S$  называется *идемпотентной*, если для любого  $x \in S$  верно, что  $x \circ x = x$ .

**Пример 1.1.4.** Рассмотрим несколько примеров идемпотентных операций.

- Операция объединения множеств  $\cup$  является идемпотентной: для любого множества  $S$  верно, что  $S \cup S = S$ .
- Операция сложения на целых числах не является идемпотентной.
- Операции “логическое и”  $\wedge$  и “логическое или”  $\vee$  являются идемпотентными.
- Операция “исключающее или” не является идемпотентной.

**Определение 1.1.8.** Пусть есть коммутативная бинарная операция  $\circ$  на множестве  $S$ . Говорят, что  $x \in S$  является *нейтральным элементом* по операции  $\circ$ , если для любого  $y \in S$  верно, что  $x \circ y = y \circ x = y$ . Если бинарная операция не является коммутативной, то можно определить *нейтральный слева* и *нейтральный справа* элементы по аналогии.

## 1.2 Полугруппа

**Определение 1.2.1.** Множество с заданной на нём ассоциативной бинарной операцией  $(S, \cdot : S \times S \rightarrow S)$  называется *полугруппой*. Если операция  $\cdot$  является коммутативной, то говорят о *коммутативной полугруппе*.

**Пример 1.2.1.** Приведём несколько примеров полугрупп.

- Положительные целые числа с операцией сложения являются полугруппой. Более того, коммутативной полугруппой.
- Целые числа с операцией взятия наибольшего из двух ( $\max$ ) являются полугруппой. Более того, коммутативной полугруппой.
- Множество всех строк конечной длины (без пустой строки) над фиксированным алфавитом  $\Sigma$  с операцией конкатенации является полугруппой. Так как конкатенация на строках не является коммутативной операцией, то и полугруппа не является коммутативной.

<sup>2</sup>Здесь может быть уместно вспомнить правила сложения дробей. Дробь с общим знаменателем складывать проще как раз из-за дистрибутивности справа.



## 1.3 Моноид

**Определение 1.3.1.** *Моноидом* называется полугруппа с нейтральным элементом. Если операция является коммутативной, то можно говорить о коммутативном моноиде.

**Пример 1.3.1.** Приведём примеры моноидов, построенных на основе полугрупп из предыдущего раздела.

- Неотрицательные целые числа с операцией сложения являются моноидом. Нейтральный элемент — 0.
- Целые числа, дополненные значением  $-\infty$  (“минус-бесконечность”) с операцией взятия наибольшего из двух ( $\max$ ) являются моноидом. Нейтральный элемент —  $-\infty$ .
- Множество всех строк конечной длины с пустой строкой (строка длины 0) над фиксированным алфавитом  $\Sigma$  и операцией конкатенации является моноидом. Нейтральный элемент — пустая строка.
- Квадратные неотрицательные матрицы<sup>3</sup> фиксированного размера с операцией умножения задают моноид. Нейтральный элемент — единичная матрица.

## 1.4 Группа

**Определение 1.4.1.** Непустое<sup>4</sup> множество  $G$  с заданной на нём бинарной операцией  $\circ : G \times G \rightarrow G$  называется *группой*  $(G, \circ)$ , если выполнены следующие аксиомы:

1. ассоциативность:  $\forall (a, b, c \in G): (a \circ b) \circ c = a \circ (b \circ c)$ ;
2. наличие нейтрального элемента:  $\exists e \in G \quad \forall a \in G: (e \circ a = a \circ e = a)$ ;
3. наличие обратного элемента:  $\forall a \in G \quad \exists a^{-1} \in G: (a \circ a^{-1} = a^{-1} \circ a = e)$ .

Иными словами, группа — это моноид с дополнительным требованием наличия обратных элементов.

**Определение 1.4.2.** Если операция  $\circ$  коммутативна, то говорят, что группа *Абелева*.

<sup>3</sup>Неотрицательной называется матрица, все элементы которой не меньше нуля.

<sup>4</sup>Требование непустоты здесь, как и далее, в определениях полукольца и кольца — дискуссионный вопрос.

раньше  
можно  
формули-  
ровать?

к примеру?

## 1.5 Полукольцо

**Определение 1.5.1.** Непустое множество  $R$  с двумя бинарными операциями  $\oplus: R \times R \rightarrow R$  (часто называют сложением) и  $\otimes: R \times R \rightarrow R$  (часто называют умножением) называется *полукольцом*, если выполнены следующие условия.

1.  $(R, \oplus)$  — это коммутативный моноид, нейтральный элемент которого —  $\mathbb{0}$ . Для любых  $a, b, c \in R$ :

- $(a \oplus b) \oplus c = a \oplus (b \oplus c)$
- $\mathbb{0} \oplus a = a \oplus \mathbb{0} = a$
- $a \oplus b = b \oplus a$

2.  $(R, \otimes)$  — это моноид, нейтральный элемент которого —  $\mathbb{1}$ . Для любых  $a, b, c \in R$ :

- $(a \otimes b) \otimes c = a \otimes (b \otimes c)$
- $\mathbb{1} \otimes a = a \otimes \mathbb{1} = a$

3.  $\otimes$  дистрибутивно слева и справа относительно  $\oplus$ :

- $a \otimes (b \oplus c) = (a \otimes b) \oplus (a \otimes c)$
- $(a \oplus b) \otimes c = (a \otimes c) \oplus (b \otimes c)$

4.  $\mathbb{0}$  является *аннигилятором* по умножению: для любого  $a \in R$   $\mathbb{0} \otimes a = a \otimes \mathbb{0} = 0$

Если операция  $\otimes$  коммутативна, то говорят о коммутативном полукольце. Если операция  $\oplus$  идемпотентна, то говорят об идемпотентном полукольце.

**Пример 1.5.1.** Рассмотрим пример полукольца, а заодно покажем, что левая и правая дистрибутивность могут существовать независимо для некоммутативного умножения<sup>5</sup>.

В качестве  $R$  возьмём множество множеств строк конечной длины над некоторым алфавитом  $\Sigma$ . В качестве сложения возьмём теоретико-множественное объединение:  $\oplus \equiv \cup$ . Нейтральный элемент по сложению — это пустое множество  $(\emptyset)$ . В качестве умножения возьмём конкатенацию множеств  $(\otimes \equiv \odot)$  и определим её следующим образом:

$$S_1 \odot S_2 = \{w_1 \cdot w_2 \mid w_1 \in S_1, w_2 \in S_2\}$$

---

<sup>5</sup>Хороший пример того, почему левую и правую дистрибутивность в случае некоммутативного умножения нужно проверять независимо (правда, для колец), приведён Николаем Александровичем Вавиловым в книге “Конкретная теория колец” на странице 6 [84].

, где  $\cdot$  — конкатенация строк. Нейтральным элементом по умножению будет являться множество из пустой строки:  $\{\varepsilon\}$ , где  $\varepsilon$  — обозначение для пустой строки.

Проверим, что  $(R, \cup, \odot)$  действительно полукольцо по нашему определению.

1.  $(R, \cup)$  — действительно коммутативный моноид с нейтральным элементом  $\emptyset$ . Для любых  $a, b, c \in R$  по свойствам теоретико-множественного объединения верно:

- $(a \cup b) \cup c = a \cup (b \cup c)$
- $\emptyset \cup a = a \cup \emptyset = a$
- $a \cup b = b \cup a$ .

2.  $(R, \odot)$  — действительно моноид с нейтральным элементом  $\{\varepsilon\}$ . Для любых  $a, b, c \in R$ :

- $(a \odot b) \odot c = a \odot (b \odot c)$  по определению  $\odot$
- $\{\varepsilon\} \odot a = \{\varepsilon \cdot w \mid w \in a\} = \{w \mid w \in a\} = a \odot \{\varepsilon\} = a$

Вообще говоря, сконструированный нами моноид не является коммутативным: легко проверить, например, что для любых непустых  $a, b \in R, a \neq b, a \neq \{\varepsilon\}, b \neq \{\varepsilon\}$ :  $a \cdot b \neq b \cdot a$  по причине некоммутативности конкатенации строк.

*Всё для  $a = \{c\}$ ;  $b = \{cc\}$*

3.  $\odot$  дистрибутивно слева и справа относительно  $\cup$ :

*$a \cdot b = b \cdot a = \{cccc\}$ .*

- $a \odot (b \cup c) = \{w_1 \cdot w_2 \mid w_1 \in a, w_2 \in b \cup c\} = \{w_1 \cdot w_2 \mid w_1 \in a, w_2 \in b\} \cup \{w_1 \cdot w_2 \mid w_1 \in a, w_2 \in c\} = (a \odot b) \cup (a \odot c)$
- Аналогично,  $(a \cup b) \odot c = (a \odot c) \cup (b \odot c)$

При этом, в общем случае,  $a \odot (b \cup c) \neq (b \cup c) \odot a$  из-за некоммутативности операции  $\odot$ .

*Пример.*

4.  $\emptyset$  является аннигилятором по умножению: для любого  $a \in R$  верно, что  $\emptyset \odot a = \{w_1 \cdot w_2 \mid w_1 \in \emptyset, w_2 \in a\} = \{w_1 \cdot w_2 \mid w_1 \in a, w_2 \in \emptyset\} = a \odot \emptyset = \emptyset$

## 1.6 Кольцо

**Определение 1.6.1.** Непустое множество  $R$  с двумя бинарными операциями  $\oplus: R \times R \rightarrow R$  (умножение) и  $\otimes: R \times R \rightarrow R$  (сложение) называется *кольцом*, если выполнены следующие условия.

1.  $(R, \oplus)$  — это Абелева группа, нейтральный элемент которой —  $0$ . Для любых  $a, b, c \in R$ :

- $(a \oplus b) \oplus c = a \oplus (b \oplus c)$

- $0 \oplus a = a \oplus 0 = a$

- $a \oplus b = b \oplus a$

- для любого  $a \in R$  существует  $-a \in R$ , такое что  $a + (-a) = 0$ .

не часто аб. гр.  
у нее есть обратные  
элементы

В последнем пункте кроется отличие от полукольца.

2.  $(R, \otimes)$  — это моноид, нейтральный элемент которого —  $1$ . Для любых  $a, b, c \in R$ :

- $(a \otimes b) \otimes c = a \otimes (b \otimes c)$

- $1 \otimes a = a \otimes 1 = a$

3.  $\otimes$  дистрибутивно слева и справа относительно  $\oplus$ :

- $a \otimes (b \oplus c) = (a \otimes b) \oplus (a \otimes c)$

- $(a \oplus b) \otimes c = (a \otimes c) \oplus (b \otimes c)$

Заметим, что мультипликативное свойство  $0$  (быть аннигилятором по умножению) не указывается явно, так как может быть выведено из остальных утверждений.

как? Добавить 0-во и  
с 0-вом

## 1.7 Матрицы и вектора

К определению матрицы мы подойдем структурно, так как в дальнейшем будем сопоставлять эту структуру с объектами различной природы, а значит определение матрицы через какой-либо из этих объектов (например через квадратичные формы) будет менее удобным.

Договоримся, что *алгебраическая структура* — это собирательное название для объектов вида “множество с набором операций” (например, кольцо, моноид, группа и т.д.), а соответствующее множество будем называть *носителем* этой структуры.

**Определение 1.7.1.** Предположим, что у нас есть некоторая алгебраическая структура с носителем  $S$ . Тогда *матрицей* будем называть прямоугольный массив размера  $n \times m$ ,  $n > 0$ ,  $m > 0$ , заполненный элементами из  $S$ .

Говорят, что  $n$  — это высота матрицы или количество строк в ней, а  $m$  — ширина матрицы или количество столбцов.

При доступе к элементам матрицы используются их индексы. При этом нумерация ведётся с левого верхнего угла, первым указывается строка, вторым — столбец. В нашей работе мы будем использовать “программистскую” традицию и нумеровать строки и столбцы с нуля<sup>6</sup>.

**Пример 1.7.1** (Матрица). Пусть есть моноид  $(S, \cdot)$ , где  $S$  — множества строк конечной длины над алфавитом  $\{a, b, c\}$ . Тогда можно построить, например, следующую матрицу  $2 \times 3$ .

$$M_{2 \times 3} = \begin{pmatrix} "a" & "ba" & "cb" \\ "ac" & "bab" & "b" \end{pmatrix}$$

$$M_{2 \times 3}[1, 1] = "bab"$$

К определению вектора мы также подойдём структурно.

**Определение 1.7.2.** *Вектором* будем называть матрицу, хотя бы один из размеров которой равен единице. Если единице равна высота матрицы, то это *вектор-строка*, если же единице равна ширина матрицы, то это *вектор-столбец*.

Операции над матрицами можно условно разделить на две группы:

- *структурные* — не зависящие от алгебраической структуры, над которой строилась матрица, и работающие только с её структурой;
- *алгебраические* — определение таковых опирается на свойства алгебраической структуры, над которой построена матрица.

Примерами структурных операций является *транспонирование*, *взятие подматрицы* и *взятие элемента по индексу*.

---

<sup>6</sup>В противоположность “математической” традиции нумеровать строки и столбцы с единицы. Стоит, правда, отметить, что в некоторых языках программирования (например, Fortran или COBOL) жива “математическая” традиция.

**Определение 1.7.3** (Транспонирование матрицы). Пусть дана матрица  $M_{n \times m}$ . Тогда результат её транспонирования  $M_{n \times m}^T = M'_{m \times n}$ , такая что  $M'[i, j] = M[j, i]$  для всех  $0 \leq i \leq m - 1$  и  $0 \leq j \leq n - 1$ .

**Пример 1.7.2.**

$$\begin{pmatrix} "a" & "ba" & "cb" \\ "ac" & "bab" & "b" \end{pmatrix}^T = \begin{pmatrix} "a" & "ac" \\ "ba" & "bab" \\ "cd" & "b" \end{pmatrix}$$

только индексы поменяются?

**Определение 1.7.4** (Взятие подматрицы). Пусть дана матрица  $M_{n \times m}$ . Тогда  $M_{n \times m}[i_0..i_1, j_0..j_1]$  — это такая  $M'_{i_1-i_0+1, j_1-j_0+1}$  что  $M'[i, j] = M[i_0 + i, j_0 + j]$  для всех  $0 \leq i \leq i_1 - i_0 + 1$  и  $0 \leq j \leq j_1 - j_0 + 1$ .

**Пример 1.7.3.**

$$\begin{pmatrix} "a" & "ba" & "cb" \\ "ac" & "bab" & "b" \end{pmatrix}[0..1, 1..2] = \begin{pmatrix} "ba" & "cb" \\ "bab" & "b" \end{pmatrix}$$

**Определение 1.7.5.** Взятие элемента по индексу — это частный случай взятия подматрицы, когда начало и конец “среза” совпадают:  $M[i, j] = M[i..i, j..j]$

**Пример 1.7.4.**

$$\begin{pmatrix} "a" & "ba" & "cb" \\ "ac" & "bab" & "b" \end{pmatrix}[0, 1] = "ba"$$

Из алгебраических операций над матрицами нас в дальнейшем будут интересовать поэлементные операции, скалярные операции, матричное умножение, произведение Кронекера.

**Определение 1.7.6** (Поэлементные операции). Пусть  $G = (S, \circ)$  — полугруппа,  $M_{n \times m}, N_{n \times m}$  — две матрицы одинакового размера над этой полугруппой. Тогда  $ewise(M, N, \circ) = P_{n \times m}$ , такая, что  $P[i, j] = M[i, j] \circ N[i, j]$ .

**Пример 1.7.5.** Пусть  $G$  — полугруппа строк с конкатенацией  $\cdot$ ,

$$M = \begin{pmatrix} "a" & "ba" & "cb" \\ "ac" & "bab" & "b" \end{pmatrix},$$

$$N = \begin{pmatrix} "c" & "aa" & "b" \\ "a" & "bac" & "bb" \end{pmatrix}.$$

Тогда

$$ewise(M, N, \cdot) = \begin{pmatrix} "ac" & "baaa" & "cbb" \\ "aca" & "babbaac" & "bbb" \end{pmatrix}.$$

**Определение 1.7.7** (Скалярная операция). Пусть  $G = (S, \circ)$  — полугруппа,  $M_{n \times m}$  — матрица над этой полугруппой,  $x \in S$ . Тогда  $scalar(M, x, \circ) = P_{n \times m}$ , такая, что  $P[i, j] = M[i, j] \circ x$ , а  $scalar(x, M, \circ) = P_{n \times m}$ , такая, что  $P[i, j] = x \circ M[i, j]$ .

**Пример 1.7.6.** Пусть  $G$  — полугруппа строк с конкатенацией  $\cdot$ ,  $x = "c"$ ,

$$M = \begin{pmatrix} "a" & "ba" & "cb" \\ "ac" & "bab" & "b" \end{pmatrix}.$$

Тогда

$$scalar(M, x, \cdot) = \begin{pmatrix} "ac" & "bac" & "cbc" \\ "acc" & "bab" & "bc" \end{pmatrix},$$

$$scalar(x, M, \cdot) = \begin{pmatrix} "ca" & "cba" & "ccb" \\ "cac" & "cbab" & "cb" \end{pmatrix}.$$

**Определение 1.7.8** (Матричное умножение). Пусть  $G = (S, \oplus, \otimes)$  — полукольцо,  $M_{n \times m}, N_{m \times k}$  — две матрицы над этим полукольцом. Тогда  $M \cdot N = P_{n \times k}$ , такая, что  $P[i, j] = \bigoplus_{0 \leq l < m} M[i, l] \otimes N[l, j]$ .

**Пример 1.7.7.** Пусть  $G$  — полукольцо из примера 1.5.1,

$$M = \begin{pmatrix} \{ "a" \} & \{ "a" \} \\ \{ "b" \} & \{ "b" \} \end{pmatrix},$$

$$N = \begin{pmatrix} \{ "c" \} \\ \{ "d" \} \end{pmatrix}.$$

Тогда

$$M \cdot N = \begin{pmatrix} \{ "a" \cdot "c" \} \cup \{ "a" \cdot "d" \} \\ \{ "b" \cdot "c" \} \cup \{ "b" \cdot "d" \} \end{pmatrix} = \begin{pmatrix} \{ "ac" \}, \{ "ad" \} \\ \{ "bc" \}, \{ "bd" \} \end{pmatrix}.$$

**Определение 1.7.9** (Произведение Кронекера). Пусть  $G = (S, \circ)$  — полугруппа,  $M_{m \times n}$  и  $N_{p \times q}$  — две матрицы над этой полугруппой. Тогда произведение Кронекера или тензорное произведение матриц  $M$  и  $N$  — это блочная матрица  $C$  размера  $mp \times nq$ , вычисляемая следующим образом:

$$C = A \otimes B = \begin{pmatrix} scalar(M[0, 0], N, \circ) & \cdots & scalar(M[0, n-1], N, \circ) \\ \vdots & \ddots & \vdots \\ scalar(M[m-1, 0], N, \circ) & \cdots & scalar(M[m-1, n-1], N, \circ) \end{pmatrix}$$

**Утверждение 1.7.1.** Отметим, что произведение Кронекера не является коммутативным. При этом всегда существуют две матрицы перестановок  $P$  и  $Q$  такие, что  $A \otimes B = P(B \otimes A)Q$ . Это свойство потребуется нам в дальнейшем.

**Пример 1.7.8.** Возьмём в качестве полугруппы целые числа с умножением.

$$M = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix}$$

$$N = \begin{pmatrix} 5 & 6 & 7 & 8 \\ 9 & 10 & 11 & 12 \\ 13 & 14 & 15 & 16 \end{pmatrix}$$

Тогда

$$\begin{aligned} M \otimes N &= \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} \otimes \begin{pmatrix} 5 & 6 & 7 & 8 \\ 9 & 10 & 11 & 12 \\ 13 & 14 & 15 & 16 \end{pmatrix} = \\ &= \begin{pmatrix} 1 \begin{pmatrix} 5 & 6 & 7 & 8 \\ 9 & 10 & 11 & 12 \\ 13 & 14 & 15 & 16 \end{pmatrix} & 2 \begin{pmatrix} 5 & 6 & 7 & 8 \\ 9 & 10 & 11 & 12 \\ 13 & 14 & 15 & 16 \end{pmatrix} \\ 3 \begin{pmatrix} 5 & 6 & 7 & 8 \\ 9 & 10 & 11 & 12 \\ 13 & 14 & 15 & 16 \end{pmatrix} & 4 \begin{pmatrix} 5 & 6 & 7 & 8 \\ 9 & 10 & 11 & 12 \\ 13 & 14 & 15 & 16 \end{pmatrix} \end{pmatrix} = \\ &= \left( \begin{array}{cccc|cccc} 5 & 6 & 7 & 8 & 10 & 12 & 14 & 16 \\ 9 & 10 & 11 & 12 & 18 & 20 & 22 & 24 \\ 13 & 14 & 15 & 16 & 26 & 28 & 30 & 32 \\ \hline 15 & 18 & 21 & 24 & 20 & 24 & 28 & 32 \\ 27 & 30 & 33 & 36 & 36 & 40 & 44 & 48 \\ 39 & 42 & 45 & 48 & 52 & 56 & 60 & 64 \end{array} \right) \end{aligned} \quad (1.1)$$

## 1.8 Теоретическая сложность умножения матриц

В рамках такого раздела теории сложности, как мелкозернистая сложность (fine-grained complexity) задача умножения двух матриц оказалась достаточно важной, так как через вычислительную сложность этой задачи можно оценить сложность большого класса различных задач. С примерами таких задач можно ознакомиться в работе [72]. Поэтому рассмотрим алгоритмы нахождения произведения двух матриц более подробно. Далее для простоты мы будем предполагать, что перемножаются две квадратные матрицы одинакового размера  $n \times n$ .



Для начал построим наивный алгоритм, сконструированный на основе определении произведения матриц. Такой алгоритм представлен на листинге 0. Его работу можно описать следующим образом: для каждой строки в первой матрице и для каждого столбца в второй матрице найти сумму произведений соответствующих элементов. Данная сумма будет значением соответствующей ячейки результирующей матрицы.

---

**Listing 1** Наивное перемножение матриц
 

---

```

1: function MATRIXMULT( $M_1, M_2, G = (S, \oplus, \otimes)$ )
2:    $M_3 =$  empty matrix of size  $n \times n$ 
3:   for  $i \in 0..n - 1$  do
4:     for  $j \in 0..n - 1$  do
5:       for  $k \in 0..n - 1$  do
6:          $M_3[i, j] = M_3[i, j] \oplus (M_1[i, k] \otimes M_2[k, j])$ 
7:   return  $M_3$ 
  
```

---

Сложность наивного произведения двух матриц составляет  $O(n^3)$ , что очевидно из псевдокода. Но можно ли улучшить этот алгоритм? Первый положительный ответ был опубликовал Ф. Штрассен в 1969 году [62]. Сложность предложенного им алгоритма —  $O(n^{\log_2 7}) \approx O(n^{2.81})$ . Основная идея — рекурсивное разбиение исходных матриц на блоки и вычисление их произведения с помощью только 7 умножений, а не 8.

Рассмотрим алгоритм Штрассена более подробно. Пусть  $A$  и  $B$  — две квадратные матрицы размера  $2^n \times 2^n$  над кольцом  $R = (S, \oplus, \otimes)$ . Наша задача найти матрицу  $C = A \cdot B$ .

Разделим матрицы  $A, B$  и  $C$  на четыре равные по размеру блока.

$$A = \begin{pmatrix} A_{1,1} & A_{1,2} \\ A_{2,1} & A_{2,2} \end{pmatrix}, B = \begin{pmatrix} B_{1,1} & B_{1,2} \\ B_{2,1} & B_{2,2} \end{pmatrix}, C = \begin{pmatrix} C_{1,1} & C_{1,2} \\ C_{2,1} & C_{2,2} \end{pmatrix}$$

На основе определения произведения матриц легко убедиться, что выполняются следующие равенства.

$$C_{1,1} = A_{1,1} \cdot B_{1,1} + A_{1,2} \cdot B_{2,1}$$

$$C_{1,2} = A_{1,1} \cdot B_{1,2} + A_{1,2} \cdot B_{2,2}$$

$$C_{2,1} = A_{2,1} \cdot B_{1,1} + A_{2,2} \cdot B_{2,1}$$

$$C_{2,2} = A_{2,1} \cdot B_{1,2} + A_{2,2} \cdot B_{2,2}$$

---

<sup>7</sup>Если размер умножаемых матриц не является натуральной степенью двойки, мы дополняем исходные матрицы дополнительными нулевыми строками и столбцами.

Тройной вложенный цикл  $O(n^3)$

можно: другое в векторе?

от этих "очевидно" можно убедиться  
мы можем  
увеселить  
себя пунши

Данная процедура не даёт нам ничего нового с точки зрения вычислительной сложности. Но мы можем двинуться дальше и определить следующие элементы.

$$P_1 \equiv (A_{1,1} + A_{2,2}) \cdot (B_{1,1} + B_{2,2})$$

$$P_2 \equiv (A_{2,1} + A_{2,2}) \cdot B_{1,1}$$

$$P_3 \equiv A_{1,1} \cdot (B_{1,2} - B_{2,2})$$

$$P_4 \equiv A_{2,2} \cdot (B_{2,1} - B_{1,1})$$

$$P_5 \equiv (A_{1,1} + A_{1,2}) \cdot B_{2,2}$$

$$P_6 \equiv (A_{2,1} - A_{1,1}) \cdot (B_{1,1} + B_{1,2})$$

$$P_7 \equiv (A_{1,2} - A_{2,2}) \cdot (B_{2,1} + B_{2,2})$$

Используя эти элементы мы можем выразить блоки результирующей матрицы следующим образом.

$$C_{1,1} = P_1 + P_4 - P_5 + P_7$$

$$C_{1,2} = P_3 + P_5$$

$$C_{2,1} = P_2 + P_4$$

$$C_{2,2} = P_1 - P_2 + P_3 + P_6$$

При таком способе вычисления мы получаем на одно умножение подматриц меньше, чем при наивном подходе. Это и приводит, в конечном итоге, к улучшению сложности всего алгоритма, который основывается на рекурсивном повторении проделанной выше процедуры.

Впоследствии сложность постепенно понижалась в ряде работ, таких как [47, 15, 54, 19, 20]. Было введено специальное обозначение для показателя степени в данной оценке:  $\omega$ . То есть сложность умножения матриц — это  $O(n^\omega)$ , и задача сводится к уменьшению значения  $\omega$ . В настоящее время работа над уменьшением показателя степени продолжается и сейчас уже предложены решения с  $\omega < 2.373^8$ .

---

<sup>8</sup>В данной области достаточно регулярно появляются новые результаты, дающие сравнительно небольшие, в терминах абсолютных величин, изменения. Так, в 2021 была представлена работа, улучшающая значение  $\omega$  в пятом знаке после запятой [3]. Несмотря на кажущуюся несерьёзность результата, подобные работы имеют большое теоретическое значение, так как улучшают наше понимание исходной задачи и её свойств.

Всё тем же Ф. Штрассеном ещё в 1969 году была выдвинута гипотеза о том, что для достаточно больших  $n$  существует алгоритм, который для любого сколь угодно маленького наперёд заданного  $\varepsilon$  перемножает матрицы за  $O(n^{2+\varepsilon})$ . На текущий момент ни доказательства, ни опровержения этой гипотезы не предъявлено.

Важной особенностью указанного выше направления улучшения алгоритмов является то, что оно допускает использования (и даже основывается на использовании) более богатых алгебраических структур, чем требуется для определения умножения двух матриц. Так, уже алгоритм Штрассена использует операцию вычитания, что приводит к необходимости иметь обратные элементы по сложению, а значит определять матрицы над кольцом. Хотя для исходного определения (1.7.8) достаточно более бедной структуры. При этом, часто, структуры, возникающие в прикладных задачах кольцами не являются. Примерами могут служить тропическое (или  $\{min, +\}$ ) полукольцо, играющее ключевую роль в тропической математике, или булево ( $\{\vee, \wedge\}$ ) полукольцо, возникающее, например, при работе с отношениями<sup>9</sup>. Значит, описанные выше решения не применимы и вопрос о существовании алгоритма с менее чем кубической сложностью снова актуален.

В попытках ответить на этот вопрос появились так называемые комбинаторные алгоритмы умножения матриц<sup>10</sup>. Классический результат в данной области — это алгоритм четырёх русских, предложенный В. Л. Арлазаровым, Е. А. Диницем, М. А. Кронродом и И. А. Фараджевым в 1970 году [82], позволяющий перемножить матрицы над конечным полукольцом за  $O(n^3/\log n)$ . Лучшим результатом<sup>11</sup> в настоящее время является алгоритм со сложностью

<sup>9</sup>Вообще говоря, в некоторых прикладных задачах возникают структуры, не являющиеся даже полукольцом. Предположим, что есть три различных множества  $S_1, S_2$  и  $S_3$  и две двухместные функции  $f : S_1 \times S_2 \rightarrow S_3$  и  $g : S_3 \times S_3 \rightarrow S_3$ . Этого достаточно, чтобы определить произведение двух матриц  $M_1$  и  $M_2$ , построенных из элементов множеств  $S_1$  и  $S_2$  соответственно. Результирующая матрица будет состоять из элементов  $S_3$ . Как видно, функции не являются бинарными операциями в смысле нашего определения. Несмотря на кажущуюся экзотичность, подобные структуры возникают на практике при работе с графами и учитываются, например, в стандарте GraphBLAS (<https://graphblas.github.io/>), где, кстати, называются полукольцами, что выглядит не вполне корректно.

<sup>10</sup>В противовес описанным выше, не являющимся комбинаторными. Стоит отметить, что строгое определение комбинаторных алгоритмов отсутствует, хотя этот термин и получил широкое употребление. В частности, Н. Бансал (Nikhil Bansal) и Р. Уильямс (Ryan Williams) в работе [8] дают определение комбинаторного алгоритма, но тут же замечают следующее: “We would like to give a definition of “combinatorial algorithm”, but this appears elusive. Although the term has been used in many of the cited references, nothing in the literature resembles a definition. For the purposes of this paper, let us think of a “combinatorial algorithm” simply as one that does not call an oracle for ring matrix multiplication.”. Ещё один вариант определения и его обсуждение можно найти в [22].

<sup>11</sup>В работе [22] предложен алгоритм со сложностью  $\Omega(n^{7/3}/2^{O(\sqrt{\log n})})$ , однако авто-

$\hat{O}(n^3/\log^4 n)^{12}$  [76].

Как видим, особенности алгебраических структур накладывают серьёзные ограничения на возможности конструирования алгоритмов. Отметим, что, хотя, в указанных случаях и предлагаются решения лучшие, чем наивное кубическое, они обладают принципиально разной асимптотической сложностью. в первом случае сложность оценивается полиномом, степень которого меньше третьей. Такие решения принято называть *истинно субкубическими* (truly subcubic). В то время как в случае комбинаторных алгоритмов степень полинома остается прежней, третьей, хотя сложность и уменьшается на логарифмический фактор. Такие решения принято называть *слегка субкубическими* (mildly subcubic). Естественный вопрос о существовании истинно субкубического алгоритма перемножения матриц над полукольцами (или же комбинаторного перемножения матриц) всё ещё не решён<sup>13</sup>.

---

ры утверждают, что сами не уверены в комбинаторности предложенного решения. По-видимому, полученные результаты ещё должны быть проверены сообществом.

<sup>12</sup>Нотация  $\hat{O}$  скрывает  $\text{poly}(\log \log)$  коэффициенты.

<sup>13</sup>Один из кандидатов — работа [22], однако на текущий момент предложенное в ней решение требует проверки.

## Глава 2

# Некоторые сведения из теории графов

В данном разделе мы дадим определения базовым понятиям из теории графов, рассмотрим несколько классических задач из области анализа графов и алгоритмы их решения. Кроме этого, поговорим о связи между линейной алгеброй и некоторыми задачами анализа графов. Всё это понадобится нам при последующей работе.

### 2.1 Основные определения

**Определение 2.1.1.** *Помеченный ориентированный граф  $\mathcal{G} = \langle V, E, L \rangle$ , где  $V$  — конечное множество вершин,  $E$  — конечное множество рёбер, т.ч.  $E \subseteq V \times L \times V$ ,  $L$  — конечное множество меток на рёбрах. В некоторых случаях метки называют *весами*<sup>1</sup> и тогда говорят о *взвешенном* графе.*

**Определение 2.1.2.** В случае, если для любого ребра  $(u, l, v)$  в графе также содержится ребро  $(v, l, u)$ , говорят, что граф *неориентированный*.

В дальнейшем речь будет идти о конечных ориентированных помеченных графах. Мы будем использовать термин *граф* подразумевая именно конечный ориентированный помеченный граф, если только не оговорено противное.

Также мы будем считать, что все вершины занумерованы подряд с нуля. То есть можно считать, что  $V$  — это отрезок  $[0, |V| - 1]$  неотрицательных целых чисел, где  $|V|$  — мощность множества  $V$ .

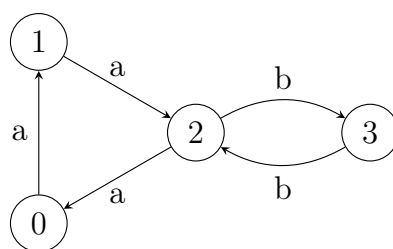
---

<sup>1</sup>Весами метки называют, как правило, тогда, когда они берутся из какого-либо поля, например  $\mathbb{R}$  или  $\mathbb{N}$ .

**Пример 2.1.1** (Пример графа и его графического представления). Пусть дан граф

$$\begin{aligned}\mathcal{G} = \langle V = \{0, 1, 2, 3\}, \\ E = \{(0, a, 1), (1, a, 2), (2, a, 0), (2, b, 3), (3, b, 2)\}, \\ L = \{a, b\}\rangle.\end{aligned}$$

Графическое представление графа  $\mathcal{G}$ :

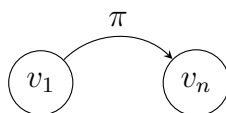


**Пример 2.1.2** (Пример рёбер графа).  $(0, a, 1)$  и  $(3, b, 2)$  — это рёбра графа  $\mathcal{G}_1$ . При этом,  $(3, b, 2)$   $(2, b, 3)$  — это разные рёбра.

**Определение 2.1.3.** Путём  $\pi$  в графе  $\mathcal{G}$  будем называть последовательность рёбер такую, что для любых двух последовательных рёбер  $e_1 = (u_1, l_1, v_1)$  и  $e_2 = (u_2, l_2, v_2)$  в этой последовательности, конечная вершина первого ребра является начальной вершиной второго, то есть  $v_1 = u_2$ . Будем обозначать путь из вершины  $v_0$  в вершину  $v_n$  как  $v_0\pi v_n$ . Иными словами,

$$v_0\pi v_n = e_0, e_1, \dots, e_{n-1} = (v_0, l_0, v_1), (v_1, l_1, v_2), \dots, (v_{n-1}, l_{n-1}, v_n).$$

Часто для представления пути мы будем использовать следующие нотации:



или

$$v_0 \xrightarrow{l_0} v_1 \xrightarrow{l_1} v_2 \xrightarrow{l_2} \dots \xrightarrow{l_{n-2}} v_{n-1} \xrightarrow{l_{n-1}} v_n.$$

**Пример 2.1.3** (Пример путей графа).  $(0, a, 1), (1, a, 2) = 0\pi_1 2$  — путь из вершины 0 в вершину 2 в графе  $\mathcal{G}_1$ . При этом,  $(0, a, 1), (1, a, 2), (2, b, 3), (3, b, 2) = 0\pi_2 2$  — это тоже путь из вершины 0 в вершину 2 в графе  $\mathcal{G}_1$ , но он не равен  $0\pi_1 2$ .

испр. сис. можно, изменив

Кроме того, нам потребуется отношение, отражающее факт существования пути между двумя вершинами.

**Определение 2.1.4.** *Отношение достижимости* в графе:  $(v_i, v_j) \in P \iff \exists v_l \pi v_j$ .

Отметим, что в некоторых задачах удобно считать по-умолчанию, что  $(v_i, v_i) \in P$ , однако наше определение такого не допускает. Исправить ситуацию можно явно добавив петли  $(v_i, l, v_i)$  для всех вершин.

Один из способов задать граф — это задать его *матрицу смежности*.

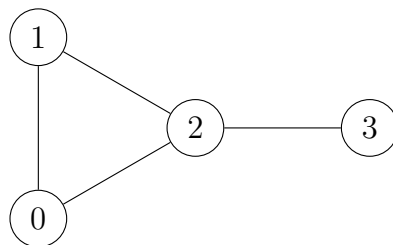
**Определение 2.1.5.** *Матрица смежности* графа  $\mathcal{G} = \langle V, E, L \rangle$  — это квадратная матрица  $M$  размера  $n \times n$ , где  $|V| = n$ , построенная над коммутативным моноидом  $\mathbb{G} = (S, \circ: S \times S \rightarrow S)$ , который конструируется следующим образом.

1.  $L \subseteq S$ .
2.  $\circ$  — коммутативная бинарная операция.
3. Существует  $\mathbb{0} \in (S \setminus L)$  — нейтральный элемент относительно  $\circ$ .

При этом  $M[i, j] = \bigcirc_{(i, l, j) \in E} l$ , где  $\bigcirc_{\emptyset} = \mathbb{0}$ .

Заметим, что наше определение матрицы смежности отличается от классического, в котором матрица является булевой и отражает лишь факт наличия хотя бы одного ребра. То есть  $M[i, j] = 1 \iff \exists e = (i, \_, j) \in E$ .

**Пример 2.1.4** (Пример матрицы смежности неориентированного графа). Пусть дан следующий неориентированный граф.



$\mathbb{G} = (S, \circ)$  в этом случае конструируется следующим образом. Во-первых, придётся предположить, что  $L$  — множество с одним элементом, скажем  $s$ , и

считать, что все рёбра помечены им<sup>2</sup>. Далее,  $S = L \cup \{n\} = \{s, n\}$ , где  $n$  — нейтральный элемент относительно  $\circ$ . Тогда  $\circ$  можно определить поточечно следующим образом.

- $s \circ s = s$
- $s \circ n = n \circ s = s$
- $n \circ n = n$

Таким образом, матрица смежности данного графа выглядит следующим образом:

$$\begin{pmatrix} n & s & s & n \\ s & n & s & n \\ s & s & n & s \\ n & n & s & n \end{pmatrix}$$

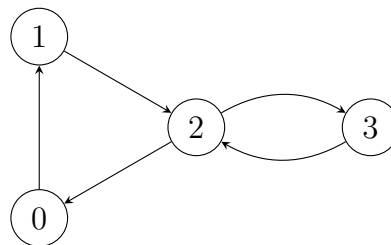
, что может показаться несколько непривычным. Однако заметим, что построенная нами структура  $\mathbb{G} = (\{s, n\}, \circ)$  изоморфна  $\mathbb{G}' = (\{1, 0\}, \vee)$ . При переходе к  $\mathbb{G}'$  мы получим привычную нам булеву матрицу смежности:

$$\begin{pmatrix} 0 & 1 & 1 & 0 \\ 1 & 0 & 1 & 0 \\ 1 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}$$

ну и зачем тогда?

Заметим, что матрица смежности неориентированного графа всегда симметрична относительно главной диагонали.

**Пример 2.1.5** (Пример матрицы смежности ориентированного графа). Дан ориентированный граф:



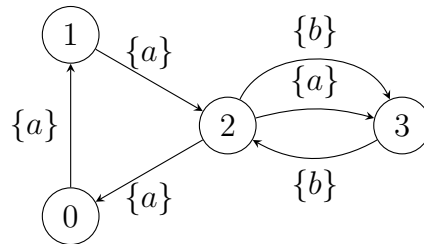
<sup>2</sup>А раз все рёбра имеют одинаковый заранее известный вес, то можно его и не писать для каждого ребра при задании графа. Поэтому привычное нам изображение получается достаточно логичным.



Построить его булеву матрицу смежности можно применив рассуждения из предыдущего примера (2.1.4) и выглядеть она будет следующим образом:

$$\begin{pmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 1 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}$$

**Пример 2.1.6** (Пример матрицы смежности помеченного графа). Пусть дан следующий помеченный граф.

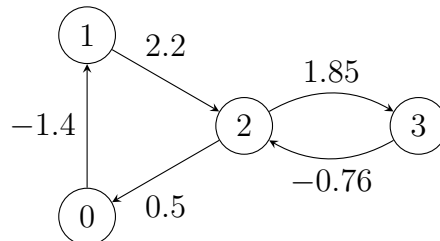


В данном случае  $L = \{\{a\}, \{b\}\}$ , а  $\mathbb{G} = (\{\{a\}, \{b\}, \{a, b\}, \emptyset\}, \cup)$ , где  $\emptyset$  — нейтральный элемент. Тогда матрица смежности исходного графа выглядит следующим образом:

$$\begin{pmatrix} \emptyset & \{a\} & \emptyset & \emptyset \\ \emptyset & \emptyset & \{a\} & \emptyset \\ \{a\} & \emptyset & \emptyset & \{a, b\} \\ \emptyset & \emptyset & \{b\} & \emptyset \end{pmatrix}$$

Необходимо заметить, что свойства  $\mathbb{G}$ , а значит и детали её построения, зависят от задачи, в рамках которой рассматривается граф. В примерах выше мы строили  $\mathbb{G}$  из некоторых общих соображений, не специфицируя решаемую задачу, стараясь получить ожидаемый результат. Далее мы рассмотрим пример, в котором видно, как решаемая задача влияет на построение  $\mathbb{G}$ .

**Пример 2.1.7** (Пример матрицы смежности взвешенного графа). Пусть дан следующий взвешенный граф:



Будем считать, что веса берутся из  $\mathbb{R}$ , а решаемая задача — поиск кратчайших путей между вершинами. В таком случае естественно предположить, что для любой вершины  $v_i$  существует петля  $(v_i, 0, v_i)$ , хотя она явно и не изображена. Далее,  $\mathbb{G} = (\mathbb{R} \cup \{\infty\}, \min)$ , где  $\infty$  — нейтральный элемент относительно операции  $\min$ .

В результате мы получим следующую матрицу смежности:

$$\begin{pmatrix} 0 & -1.4 & \infty & \infty \\ \infty & 0 & 2.2 & \infty \\ 0.5 & \infty & 0 & 1.85 \\ \infty & \infty & -0.76 & 0 \end{pmatrix}$$

Таким образом, уже можно заметить, что введение моноида как абстракции позволяет достаточно унифицированным образом смотреть на различные графы и их матрицы смежности. Далее мы увидим, что данный путь позволит решать унифицированным образом достаточно широкий круг задач, связанных с анализом путей в графах. Но сперва мы сформулируем различные варианты задачи поиска путей в графе. Это необходимо для формулировки задач, на решение которых мы в конечном итоге нацелены.

## 2.2 Задачи поиска путей

Одна из классических задач анализа графов — это задача поиска путей между вершинами с различными ограничениями.

При этом, возможны различные постановки задачи. С одной стороны, постановки различаются тем, что именно мы хотим получить в качестве результата. Здесь наиболее частыми являются следующие варианты.

- Наличие хотя бы одного пути, удовлетворяющего ограничению, в графе. В данном случае не важно, между какими вершинами существует путь, важно лишь наличие его в графе.
- Наличие пути, удовлетворяющего ограничению, между некоторыми вершинами: задача достижимости. При данной постановке задачи, нас интересует ответ на вопрос достижимости вершины  $v_i$  из вершины  $v_j$  по пути, удовлетворяющему ограничению. Такая постановка требует лишь проверить существование пути, но не его предоставления в явном виде.
- Поиск одного пути, удовлетворяющего ограничению: необходимо не только установить факт наличия пути, но и предъявить его. При этом часто

подразумевается, что возвращается любой путь, являющийся решением, без каких-либо дополнительных ограничений. Хотя, например, в некоторых задачах дополнительное требование простоты или наименьшей длины выглядит достаточно естественным.

- Поиск всех путей: необходимо предоставить все пути, удовлетворяющие заданным ограничениям.

С другой стороны, задачи могут различаться ещё и тем, как фиксируются множества стартовых и конечных вершин. Здесь возможны следующие варианты:

- от одной вершины до всех,
- между всеми парами вершин,
- между фиксированной парой вершин,
- между двумя множествами вершин  $V_1$  и  $V_2$ , что подразумевает решение задачи для всех  $(v_i, v_j) \in V_1 \times V_2$ .

Стоит отметить, что последний вариант является самым общим и остальные — лишь его частные случаи. Однако этот вариант часто выделяют отдельно, подразумевая, что остальные, выделенные, варианты в него не включаются.

В итоге, перебирая возможные варианты желаемого результата и способы фиксации стартовых и финальных вершин, мы можем сформулировать достаточно большое количество задач. Например, задачу поиска всех путей между двумя заданными вершинами, задачу поиска одного пути от фиксированной стартовой вершины до каждой вершины в графе, или задачу достижимости между всеми парами вершин.

Часто поиск путей сопровождается изучением их свойств, что далее приводит к формулированию дополнительных ограничений на пути в терминах этих свойств. Например, можно потребовать, чтобы пути были простыми или не проходили через определённые вершины. Один из естественных способов описывать свойства и, как следствие, задавать ограничения — это использовать ту алгебраическую структуру, из которой берутся веса рёбер графа<sup>3</sup>.

---

<sup>3</sup>На самом деле здесь наблюдается некоторая двойственность. С одной стороны, действительно, удобно считать, что свойства описываются в терминах некоторой заданной алгебраической структуры. Но, вместе с этим, структура подбирается исходя из решаемой задачи.

Предположим, что дан граф  $\mathcal{G} = \langle V, E, L \rangle$ , где  $L = \langle S, \oplus, \otimes \rangle$  — это полукольцо. Тогда изучение свойств путей можно описать следующим образом:

$$\{(v_i, v_j, c) \mid \exists v_i \pi v_j, c = \bigoplus_{\forall v_i \pi v_j (u, l, v) \in \pi} l\}. \quad (2.1)$$

Иными словами, для каждой пары вершин, для которой существует хотя бы один путь, их соединяющий, мы агрегируем (с помощью операции  $\oplus$  из полукольца) информацию обо всех путях между этими вершинами. При этом информация о пути получается как свёртка меток рёбер пути с использованием операции  $\otimes$ <sup>4</sup>.

*Пример добавить!*

Естественным требованием (хотя бы для прикладных задач, решаемых таким способом) является существование и конечность указанной суммы. На данном этапе мы не будем касаться того, какие именно свойства полукольца могут нам обеспечить данное свойство, однако в дальнейшем будем считать, что оно выполняется. Более того, будем стараться приводить частные для конкретной задачи рассуждения, показывающие, почему это свойство выполняется в рассматриваемых в задаче ограничениях.

Описанная выше задача общего вида называется анализом свойств путей алгебраическими методами (Algebraic Path Problem [9]) и предоставляет общий способ для решения широкого класса прикладных задач<sup>5</sup>. Наиболее известными являются такие задачи, как построение транзитивного замыкания графа и поиск кратчайших путей (All PAirs Shortest Path или APSP). Далее мы подробнее обсудим эти две задачи и предложим алгоритмы их решения.

## 2.3 Алгоритм Флойда-Уоршелла

Наиболее естественным образом решение обсуждаемых выражается в терминах операций над матрицей смежности исходного графа. Поэтому предположим, что исходный граф задан матрицей над моноидом  $\mathbb{G} = (S, \oplus)$ .

<sup>4</sup>Заметим, что детали свёртки вдоль пути зависят от свойств полукольца (и от решаемой задачи). Так, если полукольцо коммутативно, то нам не обязательно соблюдать порядок рёбер. В дальнейшем мы увидим, что данные особенности полукольца существенно влияют на особенности алгоритмов решения соответствующих задач.

<sup>5</sup>В работе “Path Problems in Networks” [9], упомянутой чуть выше, собран действительно большой список прикладных задач с описанием соответствующих полуколец. Сводная таблица на страницах 58–59 содержит 29 различных прикладных задач и соответствующих полуколец.

Как мы видели ранее, операция  $\oplus$  позволяет нам агрегировать информацию по всем параллельным рёбрам. Ровно она же и будет агрегировать информацию по всем путям между двумя вершинами<sup>6</sup>. Таким образом, осталось сконструировать операцию, отвечающую за агрегацию информации вдоль пути. Здесь мы будем исходить из того, что новый путь может быть получен из двух подпутей, а свойство нового пути зависит только от свойств исходных подпутей.

Таким образом, дополнительная операция, обозначим её  $\otimes : S \times S \rightarrow S^7$ , должна вести себя следующим образом. Пусть  $S$  — носитель моноида,  $\mathbb{0} \in S$  — нейтральный элемент относительно  $\oplus$ .

- $s_1 \otimes s_2 = s_3, s \in S, s \neq \mathbb{0}$  : если существует путь  $i\pi j$  со свойством  $s_1$  и путь  $j\pi k$  со свойством  $s_2$ , то существует путь  $i\pi k$  со свойством  $s_3$ .
- $s \otimes \mathbb{0} = \mathbb{0}$  : если существует путь  $i\pi j$  со свойством  $s$  и не существует пути  $j\pi k$ , то не существует и пути  $i\pi k$ .
- $\mathbb{0} \otimes s = \mathbb{0}$  : если не существует пути  $i\pi j$  и существует путь  $j\pi k$  со свойством  $s$ , то не существует и пути  $i\pi k$ .
- $\mathbb{0} \otimes \mathbb{0} = \mathbb{0}$  : если не существует пути  $i\pi j$  и не существует пути  $j\pi k$ , то не существует и пути  $i\pi k$ .

Зачем?

Новую операцию добавим к моноиду и получим новую алгебраическую структуру  $\mathbb{G}' = (S, \oplus, \otimes)$ . Данная структура является коммутативным моноидом по сложению (по построению) с нейтральным элементом  $\mathbb{0}$ . Из построения  $\otimes$  видно, что  $\mathbb{0}$  является аннигилятором. Ничего более про операцию *otimes*, а значит и про  $\mathbb{G}'$  мы сказать, исходя из построения, не можем. Классический фреймворк для решения algebraic path problem подразумевает, что  $\mathbb{G}'$  является полукольцом, однако далее мы увидим, что существуют задачи, в

<sup>6</sup>Вообще говоря, работая с матрицей смежности мы не видим разницу между путём и ребром, так как любая запись в матрице смежности в ячейке  $[i, j]$  говорит нам только о том, что вершины  $i$  и  $j$  связаны и эта связь обладает некоторым свойством (значение в ячейке), и ничего не говорит о том, как эта связь устроена.

<sup>7</sup>При первом рассмотрении такой выбор кажется контринтуитивным. Действительно, ведь при соединении путей мы как бы “складываем” их веса. Но при более детальном анализе поведения этой операции, в частности, относительно нейтрального элемента, становится понятно, что она ведёт себя очень похоже на умножение. Вероятно, стоит обратить внимание на операцию конкатенации, которая, с одной стороны, “делает то, что нам нужно”, а с другой, (и неспроста) часто обозначается  $\cdot$ .

которых *otimes*, например, не является ассоциативной<sup>8</sup>. А значит, согласно нашему определению,  $\mathbb{G}'$  полукольцом не является.

Теперь, когда построена алгебраическая структура, обеспечивающая вычисление формулы 2.1, мы можем предложить алгоритм вычисления этой формулы и данным алгоритмом в интересующих нас частных случаях будет являться алгоритм Флойда-Уоршелла [27, 53, 71]. Псевдокод алгоритма в обобщённом виде представлен на листинге 2, а его сложность  $O(n^3)$ . Он практически дословно основан на описанной выше идее сборки путей из двух подпутей: тройной вложенный цикл перебирает все возможные разбиения пути на две части, а в строке 7 как раз и происходит вычисление формулы 2.1.

Необходимо обратить внимание на несколько вещей. Первая — порядок обхода. Внешний цикл перебирает возможные точки разбиения (хотя мог бы, например, перебирать начальные вершины) для того, чтобы гарантировать правильный порядок вычисления подпутей (информация ни о каких подпутях не будет получена после того, как они поучаствовали в построении более длинного пути). Вторая — количество итераций. В данном случае мы ограничились тройным вложенным циклом от 0 до  $n$  и для наших задач этого будет достаточно, однако, как доказательство этого факта, так и построение аналогичного алгоритма для других задач требует аккуратного анализа решаемой задачи и последующего доказательства корректности построенного алгоритма.

---

**Listing 2** Алгоритм Флойда-Уоршелла

---

```

1: function FLOYDWARSHALL( $\mathcal{G}$ )
2:    $M \leftarrow$  матрица смежности  $\mathcal{G}$                                  $\triangleright$  Матрица над  $\mathbb{G} = (S, \oplus, \otimes)$ 
3:    $n \leftarrow |V(\mathcal{G})|$ 
4:   for  $k = 0$ ;  $k < n$ ;  $k++$  do
5:     for  $i = 0$ ;  $i < n$ ;  $i++$  do
6:       for  $j = 0$ ;  $j < n$ ;  $j++$  do
7:          $M[i, j] \leftarrow M[i, j] \oplus (M[i, k] \otimes M[k, j])$ 
8:   return  $M$ 

```

---

Хотя изначально данный алгоритм был предложен для решения задачи о кратчайших путях, при абстрагировании алгебраической структуры он превращается в алгоритм решения целого ряда задач. В частности — нахождения транзитивного замыкания. Так, если возьмём тропическое полукольцо  $(\mathbb{R}_{+\infty}, \min, +)$ , то получим алгоритм для поиска кратчайших путей. Если же

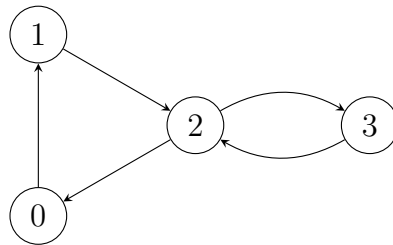
---

<sup>8</sup>Такой будет рассматриваемая в данной работе задача достижимости с ограничениями в терминах формальных языков. Другие примеры можно найти в уже упоминавшейся работе [9]

абстрагировать.

возьмём булево полукольцо, то получим алгоритм для построения транзитивного замыкания графа.

**Пример 2.3.1** (Транзитивное замыкание графа). Пусть дан следующий граф:



Его матрица смежности:

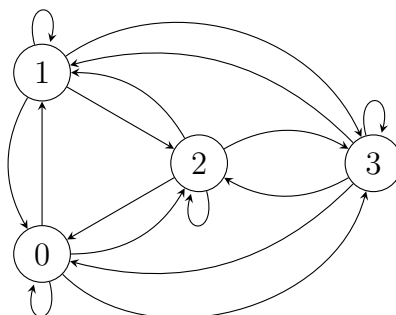
$$M = \begin{pmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 1 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}$$

Здесь мы считаем, что отношение достижимости не рефлексивно: все диагональные элементы матрицы  $M$  равны 0.

Воспользовавшись алгоритмом из листинга 2, специализированного на случай булева полукольца, можно получить следующую матрицу смежности.

$$M' = \begin{pmatrix} 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \end{pmatrix}$$

А значит, транзитивным замыканием исходного графа является полный граф и он выглядит следующим образом.



Заметим, что рефлексивность отношения (значения элементов на главной диагонали матрицы смежности) непосредственно связана с особенностями решаемой задачи. Например, если говорить о кратчайших расстояниях, то кажется естественным считать расстояние от вершины до самой себя равной нулю. Однако для задачи транзитивного замыкания это уже не столь естественное предположение и часто отдельно говорят о рефлексивно-транзитивном замыкании, которое отдельно явным образом привносит рефлексивность (петли, диагональные элементы матрицы смежности).

Аналогичным образом, используя данный алгоритм, но уже для тропического полукольца, можно решить задачу о поиске кратчайших путей для графа из примера 2.1.7.

Идеи, заложенные в алгоритме Флойда-Уоршелла, а также возможность абстрагировать его, помогут нам в дальнейшем предложить алгоритм для задачи достижимости с ограничениями в терминах формальных языков.

## 2.4 Анализ путей в графе и линейная алгебра

В данной главе мы рассмотрим некоторые связи<sup>9</sup> между графами и операциями над ними и матрицами и операциями над матрицами.

Как мы видели, достаточно естественное представление графа — это его матрица смежности. Далее можно заметить некоторое сходство между определением матричного умножения 1.7.8 и мыслями, которыми мы руководствовались, вводя операцию  $\otimes$  (2.3). Действительно, пусть есть матрица  $M$  размера  $n \times n$  над  $\mathbf{G} = (S, \oplus, \otimes)$ <sup>10</sup> — матрица смежности графа. Умножим её саму на себя: вычислим  $M' = M \cdot M$ . Тогда, по 1.7.8,  $M'[i, j] = \bigoplus_{0 \leq l < n} M[i, l] \otimes M[l, j]$ . Размер  $M'$  также  $n \times n$ . То есть  $M'$  задаёт такой граф, что в нём будет путь со свойством, являющимся агрегацией свойств всех путей, составленных из двух подпутей в исходном графе. А именно, если в исходном графе есть путь из  $i$  в  $l$  с некоторым свойством (его значение

<sup>9</sup>Связь между графами и линейной алгеброй — обширная область, в которой можно даже выделить отдельные направления, такие как спектральная теория графов. С точки зрения практики данная связь также подмечена давно и более полно с ней можно ознакомиться, например, в работах [39, 23]. Кроме этого, много полезной информации можно найти на сайте <https://graphblas.github.io/GraphBLAS-Pointers/>.

<sup>10</sup>Здесь мы уже сталкиваемся с тем, что могут иметь смысл относительно экзотические алгебраические структуры. Действительно, как мы выяснили, матрица смежности может быть определена на чем-то более бедном, чем полукольцо, а матричное умножение мы определяли над полукольцом. Но если задуматься, то и для определения произведения матриц полукольцо вовсе необязательно, достаточно более бедной структуры.



храниться в  $M[i, l]$ ), и был путь из  $l$  в  $j$  (его значение храниться в  $M[l, j]$ ), то в новом графе свойство  $M[i, l] \otimes M[l, j]$  аддитивно (используя  $\oplus$ ) учтётся в свойстве пути из  $i$  в  $j$ .

Таким образом, произведение матриц смежности соответствует обработке информации о путях интересующим нас образом. Это наблюдение позволяет предложить решение задач анализа свойств путей, основанное на операциях над матрицами. Рассмотрим такое решение для задачи о кратчайших путях.

Пусть  $D_k$  — матрица кратчайших путей, состоящих не более чем из  $k$  рёбер. То есть  $D_k[i, j]$  — это длина кратчайшего пути из вершины  $i$  в вершину  $j$ , такого, что он состоит не более чем из  $k$  ребер. Если такого пути нет, то  $D_k[i, j] = \infty$ .

Таким образом,  $D_1 = M$ , где  $M$  — это матрица смежности исходного графа, а решением APSP является  $D_{n-1}$ , вычисляемая с помощью следующего рекуррентного соотношения:

$$\begin{aligned} D(1) &= M \\ D(2) &= D(1) \cdot M = M^2 \\ D(3) &= D(2) \cdot M = M^3 \\ &\dots \\ D(n-1) &= D(n-2) \cdot M = M^{(n-1)} \end{aligned}$$

*можно  
вместе?*

Здесь мы пользуемся той особенностью задачи, что кратчайший путь в ориентированном графе (без отрицательных циклов) не может быть длиннее  $n^{11}$ . Более того, мы пользуемся тем, что используемая структура именно полукольцо, а исходное отношение рефлексивно<sup>12</sup>.

Таким образом, решение APSP сведено к произведению матриц над тропическим полукольцом, однако вычислительная сложность решения слишком большая:  $O(nK(n))$ , где  $K(n)$  — сложность алгоритма умножения матриц.

Чтобы улучшить сложность, заметим, что, например,  $D_3$  вычислять не обязательно, так как можно сразу вычислить  $D_4$  как  $D_2 \cdot D_2$ .

В итоге получим следующую последовательность вычислений:

<sup>11</sup>Раз отрицательных циклов нету, то проходить через одну вершину, коих  $n$ , больше одного раза бессмысленно.

<sup>12</sup>Тот факт, что в любой вершине есть петля с весом 0, а 0 — нейтральный для  $\otimes$ , позволяет не суммировать матрицы. Итоговая матрица содержит данные о путях длины *не больше* чем вычисляемая степень, хотя должна бы содержать данные о путях ровно и только такой длины. Предлагается самостоятельно исследовать данный феномен.

*можно  
вместе?*

$$\begin{aligned}
 D_1 &= M \\
 D_2 &= M^2 = M \cdot M \\
 D_4 &= M^4 = M^2 \cdot M^2 \\
 D_8 &= M^8 = M^4 \cdot M^4
 \end{aligned}$$

$$\begin{aligned}
 D_{2^{\log(n-1)}} &= M^{2^{\log(n-1)}} = \overset{\dots}{M^{2^{\log(n-1)-1}} \cdot M^{2^{\log(n-1)-1}}} \\
 D_{n-1} &= D_{2^{\log(n-1)}}
 \end{aligned}$$

Теперь вместо  $n$  итераций нам нужно  $\log n$ , а итоговая сложность решения —  $O(\log n K(n))$ <sup>13</sup>. Данный алгоритм называется *repeated squaring*<sup>14</sup>. Здесь мы предполагаем, что  $n-1 = 2^k$  для какого-то  $k$ . На практике такое верно далеко не всегда. Если это условие не выполняется, то необходимо взять ближайшую сверху степень двойки<sup>15</sup>.

Таким образом, APSP сводится к умножению матриц над полукольцом, что, к сожалению, не позволяет этим путём получить истинно субкубический алгоритм для задачи. тем не менее, позволяет получить слегка субкубический. Приведем некоторые работы по APSP для ориентированных графов с вещественными весами (здесь  $n$  — количество вершин в графе), по которым можно более детально ознакомиться как с историей вопроса, так и с текущими результатами:

- M.L. Fredman (1976) —  $O(n^3(\log \log n / \log n)^{\frac{1}{3}})$  — использование дерева решений [28]
- W. Dobosiewicz (1990) —  $O(n^3/\sqrt{\log n})$  — использование операций на Word Random Access Machine [24]
- T. Takaoka (1992) —  $O(n^3\sqrt{\log \log n / \log n})$  — использование таблицы поиска [63]
- Y. Han (2004) —  $O(n^3(\log \log n / \log n)^{\frac{5}{7}})$  [30]
- T. Takaoka (2004) —  $O(n^3(\log \log n)^2 / \log n)$  [64]
- T. Takaoka (2005) —  $O(n^3 \log \log n / \log n)$  [65]

<sup>13</sup>Заметим, что это оценка для худшего случая. На практике при использовании данного подхода можно прекращать вычисления как только при двух последовательных шагах получились одинаковые матрицы ( $D_i = D_{i-1}$ ). Это приводит нас к понятию *неподвижной точки*, обсуждение которого лежит за рамками повествования.

<sup>14</sup>Пример решения APSP с помощью repeated squaring: [http://users.cecs.anu.edu.au/~Alistair.Rendell/Teaching/apac\\_comp3600/module4/all\\_pairs\\_shortest\\_paths.xhtml](http://users.cecs.anu.edu.au/~Alistair.Rendell/Teaching/apac_comp3600/module4/all_pairs_shortest_paths.xhtml)

<sup>15</sup>Кажется, что это приведёт к избыточным вычислениям. Попробуйте оценить, на сколько много лишних вычислений будет сделано в худшем случае.

- U. Zwick (2004) —  $O(n^3 \sqrt{\log \log n} / \log n)$  [81]
- Т.М. Chan (2006) —  $O(n^3 / \log n)$  — многомерный принцип “разделяй и властвуй” [17]

Вопрос же о истинно субкубических алгоритмах решения APSP всё ещё открыт [18] и активно обсуждается в научном сообществе.

Аналогичным образом можно свести транзитивное замыкание к матричным операциям. Предлагаем проделать это самостоятельно, заодно обратив внимание на важность рефлексивности (в примере 2.3.1 её нет).

Заметим, что оптимизация, связанная с возведением в квадрат возможна только при ассоциативности произведения матриц, что зависит от свойств алгебраической структуры, над которой построены матрицы. Для полукольца такая оптимизация законна, однако в некоторых случаях её применить нельзя. Таким случаем как раз и будет рассматриваемая в нашей работе задача. Поэтому построение алгоритма её решения через операции над матрицами, хотя и будет основано на указанных выше соображениях, но будет сопряжено с некоторыми трудностями.



## Глава 3

# Общие сведения теории формальных языков

В данной главе мы рассмотрим основные понятия из теории формальных языков, которые пригодятся нам в дальнейшем изложении.

**Определение 3.0.1.** *Алфавит* — это конечное множество. Элементы этого множества будем называть *символами*.

**Пример 3.0.1.** Примеры алфавитов

- Латинский алфавит  $\Sigma = \{a, b, c, \dots, z\}$
- Кириллический алфавит  $\Sigma = \{а, б, в, \dots, я\}$
- Алфавит чисел в шестнадцатеричной записи

$$\Sigma = \{0, 1, 2, 3, 4, 5, 6, 7, 8, 9, A, B, C, D, E, F\}$$

Традиционное обозначение для алфавита —  $\Sigma$ . Также мы будем использовать различные прописные буквы латинского алфавита. Для обозначения символов алфавита будем использовать строчные буквы латинского алфавита:  $a, b, \dots, x, y, z$ .

Будем считать, что над алфавитом  $\Sigma$  всегда определена операция конкатенации  $(\cdot) : \Sigma^* \times \Sigma^* \rightarrow \Sigma^*$ . При записи выражений символ точки (обозначение операции конкатенации) часто будем опускать:  $a \cdot b = ab$ .

**Определение 3.0.2.** *Слово* над алфавитом  $\Sigma$  — это конечная конкатенация символов алфавита  $\Sigma$ :  $\omega = a_0 \cdot a_1 \cdot \dots \cdot a_m$ , где  $\omega$  — слово, а  $a_i \in \Sigma$  для любого  $i$ .

*св-ва операции*

**Определение 3.0.3.** Пусть  $\omega = a_0 \cdot a_1 \cdot \dots \cdot a_m$  — слово над алфавитом  $\Sigma$ . Будем называть  $m + 1$  *длиной слова* и обозначать как  $|\omega|$ .

**Определение 3.0.4.** *Язык* над алфавитом  $\Sigma$  — это множество слов над алфавитом  $\Sigma$ .

**Пример 3.0.2.** Примеры языков.

- Язык целых чисел в двоичной записи  $\{0, 1, -1, 10, 11, -10, -11, \dots\}$ .
- Язык всех правильных скобочных последовательностей

$$\{(), (()), ()(), (())(), \dots\}.$$

Любой язык над алфавитом  $\Sigma$  является подмножеством  $\Sigma^*$  — множества всех слов над алфавитом  $\Sigma$ .

Заметим, что язык не обязан быть конечным множеством, в то время как алфавит всегда конечен и изучаем мы конечные слова.

Способы задания языков

не упоминаю

- Перечислить все элементы. Такой способ работает только для конечных языков. Перечислить бесконечное множество не получится.
- Задать генератор — процедуру, которая возвращает очередное слово языка.
- Задать распознаватель — процедуру — *аю?* которая по данному слову может определить, принадлежит оно заданному языку или нет.

## Глава 4

# Задача о поиске путей с ограничениями в терминах формальных языков

В данной главе сформулируем постановку задачи о поиске путей в графе с ограничениями. А также приведём несколько примеров областей, в которых применяются алгоритмы решения этой задачи.

### 4.1 Постановка задачи

Пусть нам дан конечный ориентированный помеченный граф  $\mathcal{G} = \langle V, E, L \rangle$ . Функция  $\omega(\pi) = \omega((v_0, l_0, v_1), (v_1, l_1, v_2), \dots, (v_{n-1}, l_{n-1}, v_n)) = l_0 \cdot l_1 \cdot \dots \cdot l_{n-1}$  строит слово по пути посредством конкатенации меток рёбер вдоль этого пути. Очевидно, для пустого пути данная функция будет возвращать пустое слово, а для пути длины  $n > 0$  — непустое слово длины  $n$ .

Если теперь рассматривать задачу поиска путей, то окажется, что то множество путей, которое мы хотим найти, задаёт множество слов, то есть язык. А значит, критерий поиска мы можем сформулировать следующим образом: нас интересуют такие пути, что слова из меток вдоль них принадлежат заданному языку.

**Определение 4.1.1.** *Задача поиска путей с ограничениями в терминах формальных языков* заключается в поиске множества путей  $\Pi = \{\pi \mid \omega(\pi) \in \mathcal{L}\}$ .

В задаче поиска путей мы можем накладывать дополнительные ограничения на путь (например, чтобы он был простым, кратчайшим или Эйлеровым [40]), но это уже другая история.

Другим вариантом постановки задачи является задача достижимости.

**Определение 4.1.2.** *Задача достижимости* заключается в поиске множества пар вершин, для которых найдется путь с началом и концом в этих вершинах, что слово, составленное из меток рёбер пути, будет принадлежать заданному языку.  $\Pi' = \{(v_i, v_j) \mid \exists v_i \pi v_j, \omega(\pi) \in \mathcal{L}\}$ .

При этом, множество  $\Pi$  может являться бесконечным, тогда как  $\Pi'$  конечно, по причине конечности графа  $\mathcal{G}$ .

Язык  $\mathcal{L}$  может принадлежать разным классам и быть задан разными способами. Например, он может быть регулярным, или контекстно-свободным, или многокомпонентным контекстно-свободным.

Если  $\mathcal{L}$  — регулярный,  $\mathcal{G}$  можно рассматривать как недетерминированный конечный автомат (НКА), в котором все вершины и стартовые, и конечные. Тогда задача поиска путей, в которой  $\mathcal{L}$  — регулярный, сводится к пересечению двух регулярных языков.

Более подробно мы рассмотрим случай, когда  $\mathcal{L}$  — контекстно-свободный язык.

Путь  $G = \langle \Sigma, N, P \rangle$  — контекстно-свободная грамматика. Будем считать, что  $L \subseteq \Sigma$ . Мы не фиксируем стартовый нетерминал в определении грамматики, поэтому, чтобы описать язык, задаваемый ей, нам необходимо отдельно зафиксировать стартовый нетерминал. Таким образом, будем говорить, что  $L(G, N_i) = \{w \mid N_i \xRightarrow{*}_G w\}$  — это язык задаваемый грамматикой  $G$  со стартовым нетерминалом  $N_i$ .

**Пример 4.1.1.** Пример задачи поиска путей.

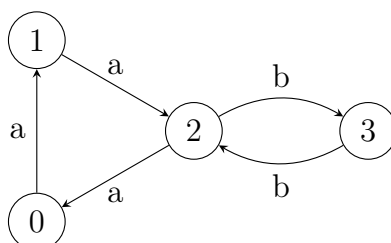
Дана грамматика  $G$  :

$$S \rightarrow ab$$

$$S \rightarrow aSb$$

Эта грамматика задаёт язык  $\mathcal{L} = a^n b^n$ .

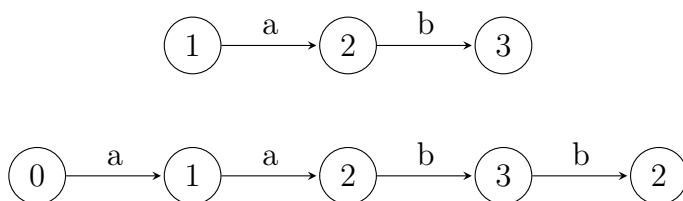
И дан граф  $\mathcal{G}$  :



не боятся  
определения  
в чуж. чуж.



*Кросс. путими ..., ева:*  
 Тогда примерами путей, принадлежащих множеству  $\Pi = \{\pi \mid \omega(\pi) \in \mathcal{L}\}$ , являются:



## 4.2 О разрешимости задачи

Задачи из определения 4.1.1 и 4.1.2 сводятся к построению пересечения языка  $\mathcal{L}$  и языка, задаваемого путями графа,  $R$ . А мы для обсуждения разрешимости задачи рассмотрим более слабую постановку задачи:

*не евр.*  
**Определение 4.2.1.** Необходимо проверить, что существует хотя бы один такой путь  $\pi$  для данного графа, для данного языка  $\mathcal{L}$ , что  $\omega(\pi) \in \mathcal{L}$ .

Эта задача сводится к проверке пустоты пересечения языка  $\mathcal{L}$  с  $R$  — регулярным языком, задаваемым графом. От класса языка  $\mathcal{L}$  зависит её разрешимость:

- Если  $\mathcal{L}$  регулярный, то получаем задачу пересечения двух регулярных языков:

$\mathcal{L} \cap R = R'$ .  $R'$  — также регулярный язык. Проверка регулярного языка на пустоту — разрешимая проблема.

- Если  $\mathcal{L}$  контекстно-свободный, то получаем задачу

$\mathcal{L} \cap R = CF$  — контекстно-свободный. Проверка контекстно-свободного языка на пустоту — разрешимая проблема.

- Помимо иерархии Хомского существуют и другие классификации языков. Так например, класс конъюнктивных (Conj) языков [45] является строгим расширением контекстно-свободных языков и все так же позволяет полиномиальный синтаксический анализ.

Пусть  $\mathcal{L}$  — конъюнктивный. При пересечении конъюнктивного и регулярного языков получается конъюнктивный ( $\mathcal{L} \cap R = Conj$ ), а проблема проверки Conj на пустоту не разрешима [46].

- Ещё один класс языков из альтернативной иерархии, не сравнимой с Иерархией Хомского, — MCFG (multiple context-free grammars) [59]. Как его частный случай — TAG (tree adjoining grammar) [36].

Если  $\mathcal{L}$  принадлежит классу MCFG, то  $\mathcal{L} \cap R$  также принадлежит MCFG. Проблема проверки пустоты MCFG разрешима [59].

Существует ещё много других классификаций языков, но поиск универсальной иерархии до сих пор продолжается.

Далее, для изучения алгоритмов решения, нас будет интересовать задача  $R \cap CF$ .

*мы будем демонстрировать на задаче...*

### 4.3 Области применения

Поиск путей с ограничениями в виде формальных языков широко применяется в различных областях. Ниже даны ключевые работы по применению поиска путей с контекстно-свободными ограничениями и ссылки на них для более детального ознакомления.

- Межпроцедурный Статанализ кода. Идея начала активно разрабатываться Томасом Репсом [51]. Далее последовал ряд, в том числе инженерных работ, применяющих достижимость с контекстно-свободными ограничениями для анализа указателей, анализа алиасов и других прикладных задач [48, 12, 79].
- Графовые БД. Впервые задача сформулирована Михалисом Яннакакисом [75]. Запросы с контекстно-свободными ограничениями нашли своё применение различных областях.
  - Социальные сети [32].
  - RDF обработка [78].
  - Биоинформатика [60].

## Глава 5

перенести  
до главы 4?

# Контекстно-свободные грамматики и языки

Из всего многообразия нас будут интересовать прежде всего контекстно-свободные грамматики.

**Определение 5.0.1.** *Контекстно-свободная грамматика* — это четвёрка вида  $\langle \Sigma, N, P, S \rangle$ , где

- $\Sigma$  — это терминальный алфавит;
- $N$  — это нетерминальный алфавит;
- $P$  — это множество правил или продукций, таких что каждая продукция имеет вид  $N_i \rightarrow \alpha$ , где  $N_i \in N$  и  $\alpha \in \{\Sigma \cup N\}^* \cup \varepsilon$ ;
- $S$  — стартовый нетерминал. Отметим, что  $\Sigma \cap N = \emptyset$ .

**Пример 5.0.1.** Грамматика, задающая язык целых чисел в двоичной записи без лидирующих нулей:  $G = \langle \{0, 1, -\}, \{S, N, A\}, P, S \rangle$ , где  $P$  определено следующим образом:

$$\begin{aligned} S &\rightarrow 0 \mid N \mid -N \\ N &\rightarrow 1A \\ A &\rightarrow 0A \mid 1A \mid \varepsilon \end{aligned}$$

При спецификации грамматики часто опускают множество терминалов и нетерминалов, оставляя только множество правил. При этом нетерминалы часто обозначаются прописными латинскими буквами, терминалы — строчными, а стартовый нетерминал обозначается буквой  $S$ . Мы будем следовать этим обозначениям, если не указано иное.

**Определение 5.0.2.** (Отношение непосредственной выводимости.) Мы говорим, что последовательность терминалов и нетерминалов  $\gamma\alpha\delta$  непосредственно выводится из  $\gamma\beta\delta$  при помощи правила  $\alpha \rightarrow \beta$  ( $\gamma\alpha\delta \Rightarrow \gamma\beta\delta$ ), если

- $\alpha \rightarrow \beta \in P$
- $\gamma, \delta \in \{\Sigma \cup N\}^* \cup \varepsilon$

**Определение 5.0.3.** Рефлексивно-транзитивное замыкание отношения — это наименьшее рефлексивное и транзитивное отношение, содержащее исходное.

**Определение 5.0.4.** Отношение выводимости является рефлексивно-транзитивным замыканием отношения непосредственной выводимости

- $\alpha \xRightarrow{*} \beta$  означает  $\exists \gamma_0, \dots, \gamma_k : \alpha \Rightarrow \gamma_0 \Rightarrow \gamma_1 \Rightarrow \dots \Rightarrow \gamma_{k-1} \Rightarrow \gamma_k \Rightarrow \beta$
- Транзитивность:  $\forall \alpha, \beta, \gamma \in \{\Sigma \cup N\}^* \cup \varepsilon : \alpha \xRightarrow{*} \beta, \beta \xRightarrow{*} \gamma \Rightarrow \alpha \xRightarrow{*} \gamma$
- Рефлексивность:  $\forall \alpha \in \{\Sigma \cup N\}^* \cup \varepsilon : \alpha \xRightarrow{*} \alpha$

*обознач*

- $\alpha \xRightarrow{*} \beta$  —  $\alpha$  выводится из  $\beta$
- $\alpha \xRightarrow{k} \beta$  —  $\alpha$  выводится из  $\beta$  за  $k$  шагов
- $\alpha \xRightarrow{+} \beta$  — при выводе использовалось хотя бы одно правило грамматики

*обознач*

**Пример 5.0.2.** Пример вывода цепочки —1101 в грамматике:

$$\begin{aligned} S &\rightarrow 0 \mid N \mid -N \\ N &\rightarrow 1A \\ A &\rightarrow 0A \mid 1A \mid \varepsilon \end{aligned}$$

$$S \Rightarrow -N \Rightarrow -1A \Rightarrow -11A \xRightarrow{*} -1101A \Rightarrow -1101$$

**Определение 5.0.5** (Вывод слова в грамматике). Слово  $\omega \in \Sigma^*$  выводимо в грамматике  $\langle \Sigma, N, P, S \rangle$ , если существует некоторый вывод этого слова из начального нетерминала  $S \xRightarrow{*} \omega$ .

**Определение 5.0.6.** Левосторонний вывод. На каждом шаге вывода заменяется самый левый нетерминал.

**Определение 5.0.7.** Правосторонний вывод. На каждом шаге вывода заменяется самый правый нетерминал.

**Пример 5.0.3.** Пример левостороннего вывода цепочки в грамматике

$$\begin{aligned} S &\rightarrow AA \mid s \\ A &\rightarrow AA \mid Bb \mid a \\ B &\rightarrow c \mid d \end{aligned}$$

$$S \Rightarrow AA \Rightarrow BbA \Rightarrow cbA \Rightarrow cbAA \Rightarrow cbaA \Rightarrow cbaa$$

Аналогично можно определить правосторонний вывод.

**Определение 5.0.8.** Язык, задаваемый грамматикой — множество строк, выводимых в грамматике  $L(G) = \{\omega \in \Sigma^* \mid S \xRightarrow{*} \omega\}$ .

**Определение 5.0.9.** Грамматики  $G_1$  и  $G_2$  называются *эквивалентными*, если они задают один и тот же язык:  $L(G_1) = L(G_2)$

**Пример 5.0.4.** Пример эквивалентных грамматик для языка целых чисел в двоичной системе счисления.

|                                                                                                                                                                                    |                                                                                                                                                              |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------|
| $\begin{aligned} \Sigma &= \{0, 1, -\} \\ N &= \{S, N, A\} \\ \\ S &\rightarrow 0 \mid N \mid -N \\ N &\rightarrow 1A \\ A &\rightarrow 0A \mid 1A \mid \varepsilon \end{aligned}$ | $\begin{aligned} \Sigma &= \{0, 1, -\} \\ N &= \{S, A\} \\ \\ S &\rightarrow 0 \mid 1A \mid -1A \\ A &\rightarrow 0A \mid 1A \mid \varepsilon \end{aligned}$ |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------|

**Определение 5.0.10.** *Неоднозначная грамматика* — грамматика, в которой существует 2 и более левосторонних (правосторонних) выводов для одного слова.

**Пример 5.0.5.** Неоднозначная грамматика для правильных скобочных последовательностей

$$S \rightarrow (S) \mid SS \mid \varepsilon$$

**Определение 5.0.11.** *Однозначная грамматика* — грамматика, в которой существует не более одного левостороннего (правостороннего) вывода для каждого слова.

**Пример 5.0.6.** Однозначная грамматика для правильных скобочных последовательностей

$$S \rightarrow (S)S \mid \varepsilon$$

**Определение 5.0.12.** *Существенно неоднозначные языки* — языки, для которых невозможно построить однозначную грамматику.

**Пример 5.0.7.** Пример существенно неоднозначного языка

$$\{a^n b^n c^m \mid n, m \in \mathbb{Z}\} \cup \{a^n b^m c^m \mid n, m \in \mathbb{Z}\}$$

сложно на 2-во

## 5.1 Дерево вывода

В некоторых случаях не достаточно знать порядок применения правил. Необходимо структурное представление вывода цепочки в грамматике. Таким представлением является *дерево вывода*.

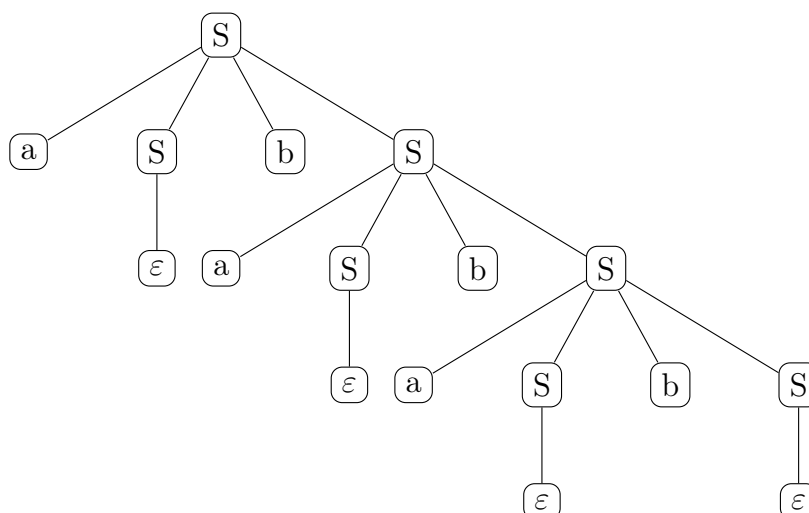
**Определение 5.1.1.** Деревом вывода цепочки  $\omega$  в грамматике  $G = \langle \Sigma, N, S, P \rangle$  называется дерево, удовлетворяющее следующим свойствам.

1. Помеченное: метка каждого внутреннего узла — нетерминал, метка каждого листа — терминал или  $\varepsilon$ .
2. Корневое: корень помечен стартовым нетерминалом.
3. Упорядоченное.
4. В дереве может существовать узел с меткой  $N_i$  и сыновьями  $M_j \dots M_k$  только тогда, когда в грамматике есть правило вида  $N_i \rightarrow M_j \dots M_k$ .
5. Крона образует исходную цепочку  $\omega$ .

**Пример 5.1.1.** Построим дерево вывода цепочки  $ababab$  в грамматике

$$G = \langle \{a, b\}, \{S\}, S, \{S \rightarrow a S b S, S \rightarrow \varepsilon\} \rangle$$

универсальная свобода  
грамматик



**Теорема 5.1.1.** Пусть  $G = \langle \Sigma, N, P, S \rangle$  — КС-грамматика. Вывод  $S \xRightarrow{*} \alpha$ , где  $\alpha \in (N \cup \Sigma)^*$ ,  $\alpha \neq \varepsilon$  существует  $\Leftrightarrow$  существует дерево вывода в грамматике  $G$  с кроной  $\alpha$ .

ссылка на 2-во

## 5.2 Пустота КС-языка

**Теорема 5.2.1.** Существует алгоритм, определяющий, является ли язык, порождаемый КС грамматикой, пустым.

*Доказательство.* Следующая лемма утверждает, что если в КС языке есть выводимое слово, то существует другое выводимое слово с деревом вывода не глубже количества нетерминалов грамматики. Для доказательства теоремы достаточно привести алгоритм, последовательно строящий все деревья глубины не больше количества нетерминалов грамматики, и проверяющий, являются ли такие деревья деревьями вывода. Если в результате работы алгоритма не удалось построить ни одного дерева, то грамматика порождает пустой язык.  $\square$

**Лемма 5.2.2.** Если в данной грамматике выводится некоторая цепочка, то существует цепочка, дерево вывода которой не содержит ветвей длиннее  $m$ , где  $m$  — количество нетерминалов грамматики.

*Доказательство.* Рассмотрим дерево вывода цепочки  $\omega$ . Если в нем есть 2 узла, соответствующих одному нетерминалу  $A$ , обозначим их  $n_1$  и  $n_2$ .

Предположим,  $n_1$  расположен ближе к корню дерева, чем  $n_2$ .

$S \xRightarrow{*} \alpha A_{n_1} \beta \xRightarrow{*} \alpha \omega_1 \beta; S \xRightarrow{*} \alpha \gamma A_{n_2} \delta \beta \xRightarrow{*} \alpha \gamma \omega_2 \delta \beta$ , при этом  $\omega_2$  является подцепочкой  $\omega_1$ .

Заменяем в изначальном дереве узел  $n_1$  на  $n_2$ . Полученное дерево является деревом вывода  $\alpha \omega_2 \delta$ .

Повторяем процесс замены одинаковых нетерминалов до тех пор, пока в дереве не останутся только уникальные нетерминалы.

В полученном дереве не может быть ветвей длины большей, чем  $m$ .

По построению оно является деревом вывода.  $\square$

### 5.3 Нормальная форма Хомского

**Определение 5.3.1.** Контекстно-свободная грамматика  $\langle \Sigma, N, P, S \rangle$  находится в *Нормальной Форме Хомского*, если она содержит только правила следующего вида:

- $A \rightarrow BC$ , где  $A, B, C \in N$ ,  $S$  не содержится в правой части правила
- $A \rightarrow a$ , где  $A \in N, a \in \Sigma$
- $S \rightarrow \varepsilon$

**Теорема 5.3.1.** Любую КС грамматику можно преобразовать в НФХ.

*Доказательство.* Алгоритм преобразования в НФХ состоит из следующих шагов:

- Замена неодинокных терминалов
- Удаление длинных правил
- Удаление  $\varepsilon$ -правил
- Удаление цепных правил
- Удаление бесполезных нетерминалов

То, что каждый из этих шагов преобразует грамматику к эквивалентной, при этом является алгоритмом, доказано в следующих леммах.  $\square$

**Лемма 5.3.2.** Для любой КС-грамматики можно построить эквивалентную, которая не содержит правила с неодинокными терминалами.



*Доказательство.* Каждое правило  $A \rightarrow B_0 B_1 \dots B_k, k \geq 1$  заменить на множество правил:

- $A \rightarrow C_0 C_1 \dots C_k$
- $\{C_i \rightarrow B_i \mid B_i \in \Sigma, C_i \text{ — новый нетерминал}\}$

□

**Лемма 5.3.3.** Для любой КС-грамматики можно построить эквивалентную, которая не содержит правил длины больше 2.

*Доказательство.* Каждое правило  $A \rightarrow B_0 B_1 \dots B_k, k \geq 2$  заменить на множество правил:

- $A \rightarrow B_0 C_0$
- $C_0 \rightarrow B_1 C_1$
- ...
- $C_{k-3} \rightarrow B_{k-2} C_{k-2}$
- $C_{k-2} \rightarrow B_{k-1} B_k$

□

**Лемма 5.3.4.** Для любой КС-грамматики можно построить эквивалентную, не содержащую  $\varepsilon$ -правил.

*Доказательство.* Определим  $\varepsilon$ -правила:

- $A \rightarrow \varepsilon$
- $A \rightarrow B_0 \dots B_k, \forall i : B_i \text{ — } \varepsilon\text{-правило.}$

Каждое правило  $A \rightarrow B_0 B_1 \dots B_k$  заменяем на множество правил, где каждое  $\varepsilon$ -правило удалено во всех возможных комбинациях.

**Лемма 5.3.5.** Можно удалить все цепные правила

*Доказательство.* *Цепное правило* — правило вида  $A \rightarrow B$ , где  $A, B \in N$ . *Цепная пара* — упорядоченная пара  $(A, B)$ , в которой  $A \xRightarrow{*} B$ , используя только цепные правила.

Алгоритм:

нужно  
убрать  
то, что  
во всех  
комб.

так вообще работает? можно  
через рекурсию?

1. Найти все цепные пары в грамматике  $G$ . Найти все цепные пары можно по индукции: Базис:  $(A, A)$  — цепная пара для любого нетерминала, так как  $A \xRightarrow{*} A$  за ноль шагов. Индукция: Если пара  $(A, B_0)$  — цепная, и есть правило  $B_0 \rightarrow B_1$ , то  $(A, B_1)$  — цепная пара.
2. Для каждой цепной пары  $(A, B)$  добавить в грамматику  $G'$  все правила вида  $A \rightarrow a$ , где  $B \rightarrow a$  — нецепное правило из  $G$ .
3. Удалить все цепные правила

Пусть  $G$  — контекстно-свободная грамматика.  $G'$  — грамматика, полученная в результате применения алгоритма к  $G$ . Тогда  $L(G) = L(G')$ .  $\square$

**Определение 5.3.2.** Нетерминал  $A$  называется *порождающим*, если из него может быть выведена конечная терминальная цепочка. Иначе он называется *непорождающим*.

**Лемма 5.3.6.** Можно удалить все бесполезные (непорождающие) нетерминалы

*Доказательство.* После удаления из грамматики правил, содержащих непорождающие нетерминалы, язык не изменится, так как непорождающие нетерминалы по определению не могли участвовать в выводе какого-либо слова.

Алгоритм нахождения порождающих нетерминалов:

1. Множество порождающих нетерминалов пустое.
2. Найти правила, не содержащие нетерминалов в правых частях и добавить нетерминалы, встречающихся в левых частях таких правил, в множество.
3. Если найдено такое правило, что все нетерминалы, стоящие в его правой части, уже входят в множество, то добавить в множество нетерминалы, стоящие в его левой части.
4. Повторить предыдущий шаг, если множество порождающих нетерминалов изменилось.

В результате получаем множество всех порождающих нетерминалов грамматики, а все нетерминалы, не попавшие в него, являются непорождающими. Их можно удалить.  $\square$

**Пример 5.3.1.** Приведем в Нормальную Форму Хомского однозначную грамматику правильных скобочных последовательностей:  $S \rightarrow aSbS \mid \varepsilon$

Первым шагом добавим новый нетерминал и сделаем его стартовым:

$$\begin{aligned} S_0 &\rightarrow S \\ S &\rightarrow aSbS \mid \varepsilon \end{aligned}$$

Заменяем все терминалы на новые нетерминалы:

$$\begin{aligned} S_0 &\rightarrow S \\ S &\rightarrow LSRS \mid \varepsilon \\ L &\rightarrow a \\ R &\rightarrow b \end{aligned}$$

Избавимся от длинных правил:

$$\begin{aligned} S_0 &\rightarrow S \\ S &\rightarrow LS' \mid \varepsilon \\ S' &\rightarrow SS'' \\ S'' &\rightarrow RS \\ L &\rightarrow a \\ R &\rightarrow b \end{aligned}$$

Избавимся от  $\varepsilon$ -продукций:

$$\begin{aligned} S_0 &\rightarrow S \mid \varepsilon \\ S &\rightarrow LS' \\ S' &\rightarrow S'' \mid SS'' \\ S'' &\rightarrow R \mid RS \\ L &\rightarrow a \\ R &\rightarrow b \end{aligned}$$

Избавимся от цепных правил:

$$\begin{aligned}
S_0 &\rightarrow LS' \mid \varepsilon \\
S &\rightarrow LS' \\
S' &\rightarrow b \mid RS \mid SS'' \\
S'' &\rightarrow b \mid RS \\
L &\rightarrow a \\
R &\rightarrow b
\end{aligned}$$

**Определение 5.3.3.** Контекстно-свободная грамматика  $\langle \Sigma, N, P, S \rangle$  находится в *ослабленной Нормальной Форме Хомского*, если она содержит только правила следующего вида:

- $A \rightarrow BC$ , где  $A, B, C \in N$
- $A \rightarrow a$ , где  $A \in N, a \in \Sigma$
- $A \rightarrow \varepsilon$ , где  $A \in N$

То есть ослабленная НФХ отличается от НФХ тем, что:

1.  $\varepsilon$  может выводиться из любого нетерминала
2.  $S$  может появляться в правых частях правил

*Зачем это определение? Скорее*

## 5.4 Лемма о накачке

**Лемма 5.4.1.** Пусть  $L$  — контекстно-свободный язык над алфавитом  $\Sigma$ , тогда существует такое  $n$ , что для любого слова  $\omega \in L$ ,  $|\omega| \geq n$  найдутся слова  $u, v, x, y, z \in \Sigma^*$ , для которых верно:  $uvxyz = \omega$ ,  $vy \neq \varepsilon$ ,  $|vxy| \leq n$  и для любого  $k \geq 0$   $uv^kxy^kz \in L$ .

Идея доказательства леммы о накачке.

1. Для любого КС языка можно найти грамматику в нормальной форме Хомского.
2. Очевидно, что если брать достаточно длинные цепочки, то в дереве вывода этих цепочек, на пути от корня к какому-то листу обязательно будет нетерминал, встречающийся минимум два раза. Если  $m$  — количество нетерминалов в НФХ, то длины  $2^{m+1}$  должно хватить. Это и будет  $n$  из леммы.

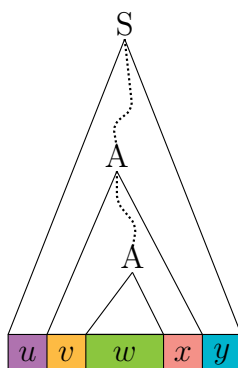


Рис. 5.1: Разбиение цепочки для леммы о накачке

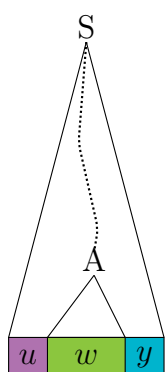
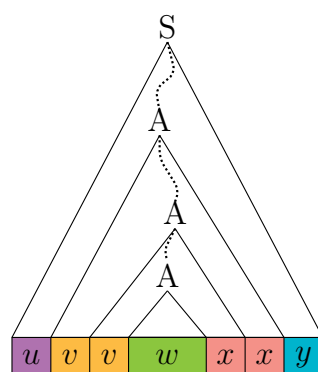
(a)  $k = 0$ .(b)  $k = 2$ .

Рис. 5.2: Пример накачки цепочки с рисунка 5.1

3. Возьмём путь, на котором есть хотя бы дважды повторяется некоторый нетерминал. Скажем, это нетерминал  $N_1$ . Пойдём от листа по этому пути. Найдём первое появление  $N_1$ . Цепочка, задаваемая поддеревом для этого узла — это  $x$  из леммы.
4. Пойдём дальше и найдём второе появление  $N_1$ . Цепочка, задаваемая поддеревом для этого узла — это  $vxy$  из леммы.
5. Теперь мы можем копировать кусок дерева между этими повторениями  $N_1$  и таким образом накачивать исходную цепочку.

Надо только проверить выполнение ограничений на длины.

Нахождение разбиения и пример накачки продемонстрированы на рисунках 5.1 и 5.2.

Для примера предлагается проверить неконтекстность языка  $L = \{a^n b^n \mid n > 0\}$ .

*надо это проверить*

## 5.5 Замкнутость КС языков относительно операций

**Теорема 5.5.1.** Контекстно-свободные языки замкнуты относительно следующих операций:

1. Объединение: если  $L_1$  и  $L_2$  — контекстно-свободные языки, то и  $L_3 = L_1 \cup L_2$  — контекстно-свободный.
2. Конкатенация: если  $L_1$  и  $L_2$  — контекстно-свободные языки, то и  $L_3 = L_1 \cdot L_2$  — контекстно-свободный.
3. Замыкание Клини: если  $L_1$  — контекстно-свободный, то и  $L_2 = \bigcup_{i=0}^{\infty} L_1^i$  — контекстно-свободный.
4. Разворот: если  $L_1$  — контекстно-свободный, то и  $L_2 = L_1^r$  — контекстно-свободный.
5. Пересечение с регулярными языками: если  $L_1$  — контекстно-свободный, а  $L_2$  — регулярный, то  $L_3 = L_1 \cap L_2$  — контекстно-свободный.
6. Разность с регулярными языками: если  $L_1$  — контекстно-свободный, а  $L_2$  — регулярный, то  $L_3 = L_1 \setminus L_2$  — контекстно-свободный.

Для доказательства пунктов 1–4 можно построить КС грамматику нового языка имея грамматики для исходных. Будем предполагать, что множества нетерминальных символов различных грамматик для исходных языков не пересекаются.

1.  $G_1 = \langle \Sigma_1, N_1, P_1, S_1 \rangle$  — грамматика для  $L_1$ ,  $G_2 = \langle \Sigma_2, N_2, P_2, S_2 \rangle$  — грамматика для  $L_2$ , тогда  $G_3 = \langle \Sigma_1 \cup \Sigma_2, N_1 \cup N_2 \cup \{S_3\}, P_1 \cup P_2 \cup \{S_3 \rightarrow S_1 \mid S_2\}, S_3 \rangle$  — грамматика для  $L_3$ .
2.  $G_1 = \langle \Sigma_1, N_1, P_1, S_1 \rangle$  — грамматика для  $L_1$ ,  $G_2 = \langle \Sigma_2, N_2, P_2, S_2 \rangle$  — грамматика для  $L_2$ , тогда  $G_3 = \langle \Sigma_1 \cup \Sigma_2, N_1 \cup N_2 \cup \{S_3\}, P_1 \cup P_2 \cup \{S_3 \rightarrow S_1 S_2\}, S_3 \rangle$  — грамматика для  $L_3$ .

3.  $G_1 = \langle \Sigma_1, N_1, P_1, S_1 \rangle$  — грамматика для  $L_1$ , тогда  $G_2 = \langle \Sigma_1, N_1 \cup \{S_2\}, P_1 \cup \{S_2 \rightarrow S_1 S_2 \mid \varepsilon\}, S_2 \rangle$  — грамматика для  $L_2$ .
4.  $G_1 = \langle \Sigma_1, N_1, P_1, S_1 \rangle$  — грамматика для  $L_1$ , тогда  $G_2 = \langle \Sigma_1, N_1, \{N^i \rightarrow \omega^R \mid N^i \rightarrow \omega \in P_1\}, S_1 \rangle$  — грамматика для  $L_2$ .

*не определяется*

Чтобы доказать замкнутость относительно пересечения с регулярными языками, построим по КС грамматике рекурсивный автомат  $R_1$ , по регулярному выражению — детерминированный конечный автомат  $R_2$ , и построим их прямое произведение  $R_3$ . Переходы по терминальным символам в новом автомате возможны тогда и только тогда, когда они возможны одновременно и в исходном рекурсивном автомате и в исходном конечном. За рекурсивные вызовы отвечает исходный рекурсивный автомат. Значит цепочка принимается  $R_3$  тогда и только тогда, когда она принимается одновременно  $R_1$  и  $R_2$ : так как состояния  $R_3$  — это пары из состояния  $R_1$  и  $R_2$ , то по трассе вычислений  $R_3$  мы всегда можем построить трассу для  $R_1$  и  $R_2$  и наоборот.

Чтобы доказать замкнутость относительно разности с регулярным языком, достаточно вспомнить, что регулярные языки замкнуты относительно дополнения, и выразить разность через пересечение с дополнением:

$$L_1 \setminus L_2 = L_1 \cap \overline{L_2}$$

□

**Теорема 5.5.2.** Контекстно-свободные языки не замкнуты относительно следующих операций:

1. Пересечение: если  $L_1$  и  $L_2$  — контекстно-свободные языки, то и  $L_3 = L_1 \cap L_2$  — не контекстно-свободный. — *никогда не-КС*
2. Разность: если  $L_1$  и  $L_2$  — контекстно-свободные языки, то и  $L_3 = L_1 \setminus L_2$  — не контекстно-свободный. — ~~никогда не-КС~~

Чтобы доказать незамкнутость относительно пересечения, рассмотрим языки  $L_1 = \{a^n b^n c^k \mid n \geq 0, k \geq 0\}$  и  $L_2 = \{a^k b^n c^n \mid n \geq 0, k \geq 0\}$ . Очевидно, что  $L_1$  и  $L_2$  — контекстно-свободные языки. Рассмотрим  $L_3 = L_1 \cap L_2 = \{a^n b^n c^n \mid n \geq 0\}$ .  $L_3$  не является контекстно-свободным по лемме о накачке для контекстно-свободных языков.

Чтобы доказать незамкнутость относительно разности сделаем следующее.

1. Рассмотрим языки  $L_4 = \{a^m b^n c^k \mid m \neq n, k \geq 0\}$  и  $L_5 = \{a^m b^n c^k \mid n \neq k, m \geq 0\}$ . Эти языки являются контекстно-свободными. Это легко заметить, если знать, что язык  $L'_4 = \{a^m b^n c^k \mid 0 \leq m < n, k \geq 0\}$  задаётся следующей грамматикой:

$$\begin{array}{ll} S \rightarrow Sc & T \rightarrow aTb \\ S \rightarrow T & T \rightarrow Tb \\ & T \rightarrow b. \end{array}$$

2. Рассмотрим язык  $L_6 = \overline{L'_6} = \overline{\{a^n b^m c^k \mid n \geq 0, m \geq 0, k \geq 0\}}$ . Данный язык является регулярным. *u?u?*
3. Рассмотрим язык  $L_7 = L_4 \cup L_5 \cup L_6$  — контекстно-свободный, так как является объединением контекстно-свободных.
4. Рассмотрим  $\overline{L_7} = \{a^n b^n c^n \mid n \geq 0\} = L_3$ :  $L_4$  и  $L_5$  задают языки с правильным порядком символов, но неравным их количеством,  $L_6$  задаёт язык с неправильным порядком символов. Из предыдущего пункта мы знаем, что  $L_3$  не является контекстно-свободным. □

↑

переходимые  
языки, чтобы  
показать

показать  
что-то



## Глава 6

# СҮК для вычисления КС запросов

В данной главе мы рассмотрим алгоритм СҮК, позволяющий установить принадлежность слова грамматике и предоставить его вывод, если таковой имеется.

Наш главный интерес заключается в возможности применения данного алгоритма для решения описанной в предыдущей главе задачи — поиска путей с ограничениями в терминах формальных языков. Как уже было указано выше, будем рассматривать случай контекстно-свободных языков.

### 6.1 Алгоритм СҮК

Алгоритм СҮК (Cocke-Younger-Kasami) — один из классических алгоритмов синтаксического анализа. Его асимптотическая сложность в худшем случае —  $O(n^3 * |N|)$ , где  $n$  — длина входной строки, а  $N$  — количество нетерминалов во входной грамматике [35].

Для его применения необходимо, чтобы подаваемая на вход грамматика находилась в Нормальной Форме Хомского (НФХ) (5.3). Других ограничений нет и, следовательно, данный алгоритм применим для работы с произвольными контекстно-свободными языками.

В основе алгоритма лежит принцип динамического программирования. Используются два соображения:

1. Для правила вида  $A \rightarrow a$  очевидно, что из  $A$  выводится  $\omega$  (с применением этого правила) тогда и только тогда, когда  $a = \omega$ :

$$A \xRightarrow{*} \omega \iff \omega = a$$

не предыдущ.  
использ. ref 8  
label

?

2. Для правила вида  $A \rightarrow BC$  понятно, что из  $A$  выводится  $\omega$  (с применением этого правила) тогда и только тогда, когда существуют две цепочки  $\omega_1$  и  $\omega_2$  такие, что  $\omega_1$  выводима из  $B$ ,  $\omega_2$  выводима из  $C$  и при этом  $\omega = \omega_1\omega_2$ :

$$A \Rightarrow BC \stackrel{*}{\Rightarrow} \omega \iff \exists \omega_1, \omega_2 : \omega = \omega_1\omega_2, B \stackrel{*}{\Rightarrow} \omega_1, C \stackrel{*}{\Rightarrow} \omega_2$$

Или в терминах позиций в строке:

$$A \Rightarrow BC \stackrel{*}{\Rightarrow} \omega \iff \exists k \in [1 \dots |\omega|] : B \stackrel{*}{\Rightarrow} \omega[1 \dots k], C \stackrel{*}{\Rightarrow} \omega[k+1 \dots |\omega|]$$

В процессе работы алгоритма заполняется булева трехмерная матрица  $M$  размера  $n \times n \times |N|$  таким образом, что

$$M[i, j, A] = \text{true} \iff A \stackrel{*}{\Rightarrow} \omega[i \dots j]$$

.

Первым шагом инициализируем матрицу, заполнив значения  $M[i, i, A]$ :

- $M[i, i, A] = \text{true}$ , если в грамматике есть правило  $A \rightarrow \omega[i]$ .
- $M[i, i, A] = \text{false}$ , иначе.

Далее используем динамику: на шаге  $m > 1$  предполагаем, что ячейки матрицы  $M[i', j', A]$  заполнены для всех нетерминалов  $A$  и пар  $i', j' : j' - i' < m$ . Тогда можно заполнить ячейки матрицы  $M[i, j, A]$ , где  $j - i = m$  следующим образом:

$$M[i, j, A] = \bigvee_{A \rightarrow BC} \bigvee_{k=i}^{j-1} M[i, k, B] \wedge M[k, j, C]$$

По итогу работы алгоритма значение в ячейке  $M[0, |\omega|, S]$ , где  $S$  — стартовый нетерминал грамматики, отвечает на вопрос о выводимости цепочки  $\omega$  в грамматике.

**Пример 6.1.1.** Рассмотрим пример работы алгоритма СΥΚ на грамматике правильных скобочных последовательностей в Нормальной Форме Хомского.

$$\begin{array}{lll}
 S \rightarrow AS_2 \mid \varepsilon & S_2 \rightarrow b \mid BS_1 \mid S_1S_3 & A \rightarrow a \\
 S_1 \rightarrow AS_2 & S_3 \rightarrow b \mid BS_1 & B \rightarrow b
 \end{array}$$

Проверим выводимость цепочки  $\omega = aabbab$ .

Так как трехмерные матрицы рисовать на двумерной бумаге не очень удобно, мы будем иллюстрировать работу алгоритма двумерными матрицами размера  $n \times n$ , где в ячейках указано множество нетерминалов, выводящих соответствующую подстроку.

Шаг 1. Инициализируем матрицу элементами на главной диагонали:

$$\begin{pmatrix}
 \{A\} & \emptyset & \emptyset & \emptyset & \emptyset & \emptyset \\
 \emptyset & \{A\} & \emptyset & \emptyset & \emptyset & \emptyset \\
 \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset & \emptyset \\
 \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset \\
 \emptyset & \emptyset & \emptyset & \emptyset & \{A\} & \emptyset \\
 \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\}
 \end{pmatrix}$$

Шаг 2. Заполняем диагональ, находящуюся над главной:

$$\begin{pmatrix}
 \{A\} & \emptyset & \emptyset & \emptyset & \emptyset & \emptyset \\
 \emptyset & \{A\} & \{S_1\} & \emptyset & \emptyset & \emptyset \\
 \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset & \emptyset \\
 \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset \\
 \emptyset & \emptyset & \emptyset & \emptyset & \{A\} & \{S_1\} \\
 \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\}
 \end{pmatrix}$$

В двух ячейках появились нетерминалы  $S_1$  благодаря присутствию в грамматике правила  $S_1 \rightarrow AS_2$ .

Шаг 3. Заполняем следующую диагональ:

$$\begin{pmatrix}
 \{A\} & \emptyset & \emptyset & \emptyset & \emptyset & \emptyset \\
 \emptyset & \{A\} & \{S_1\} & \{S_2\} & \emptyset & \emptyset \\
 \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset & \emptyset \\
 \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \{S_2, S_3\} \\
 \emptyset & \emptyset & \emptyset & \emptyset & \{A\} & \{S_1\} \\
 \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\}
 \end{pmatrix}$$

Переформат,  
чтобы строка  
была в  
ячейке  
и левой  
колонке?

Шаг 4. И следующую за ней:

$$\begin{pmatrix} \{A\} & \emptyset & \emptyset & \{S_1, S\} & \emptyset & \emptyset \\ \emptyset & \{A\} & \{S_1\} & \{S_2\} & \emptyset & \emptyset \\ \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset & \emptyset \\ \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \{S_2, S_3\} \\ \emptyset & \emptyset & \emptyset & \emptyset & \{A\} & \{S_1\} \\ \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} \end{pmatrix}$$

Шаг 5 Заполняем предпоследнюю диагональ:

$$\begin{pmatrix} \{A\} & \emptyset & \emptyset & \{S_1, S\} & \emptyset & \emptyset \\ \emptyset & \{A\} & \{S_1\} & \{S_2\} & \emptyset & \{S_2\} \\ \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset & \emptyset \\ \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \{S_2, S_3\} \\ \emptyset & \emptyset & \emptyset & \emptyset & \{A\} & \{S_1\} \\ \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} \end{pmatrix}$$

Шаг 6. Завершаем выполнение алгоритма:

$$\begin{pmatrix} \{A\} & \emptyset & \emptyset & \{S_1, S\} & \emptyset & \{S_1, S\} \\ \emptyset & \{A\} & \{S_1\} & \{S_2\} & \emptyset & \{S_2\} \\ \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset & \emptyset \\ \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \{S_2, S_3\} \\ \emptyset & \emptyset & \emptyset & \emptyset & \{A\} & \{S_1\} \\ \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} \end{pmatrix}$$

Стартовый нетерминал находится в верхней правой ячейке, а значит цепочка  $aabbab$  выводима в нашей грамматике.

**Пример 6.1.2.** Теперь выполним алгоритм на цепочке  $\omega = aba a$ .

Шаг 1. Инициализируем таблицу:

$$\begin{pmatrix} \{A\} & \emptyset & \emptyset & \emptyset \\ \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset \\ \emptyset & \emptyset & \{A\} & \emptyset \\ \emptyset & \emptyset & \emptyset & \{A\} \end{pmatrix}$$

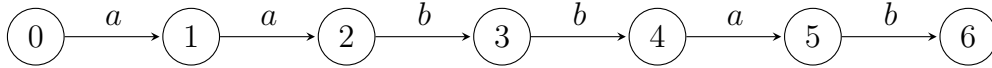
Шаг 2. Заполняем следующую диагональ:

$$\begin{pmatrix} \{A\} & \{S_1, S\} & \emptyset & \emptyset \\ \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset \\ \emptyset & \emptyset & \{A\} & \emptyset \\ \emptyset & \emptyset & \emptyset & \{A\} \end{pmatrix}$$

Больше ни одну ячейку в таблице заполнить нельзя и при этом стартовый нетерминал отсутствует в правой верхней ячейке, а значит эта строка не выводится в грамматике правильных скобочных последовательностей.

## 6.2 Алгоритм для графов на основе СΥΚ

Первым шагом на пути к решению задачи достижимости с использованием СΥΚ является модификация представления входа. Прежде мы сопоставляли каждому символу слова его позицию во входной цепочке, поэтому при инициализации заполняли главную диагональ матрицы. Теперь вместо этого обозначим числами позиции между символами. В результате слово можно представить в виде линейного графа следующим образом (в качестве примера рассмотрим слово *aabbab* из предыдущей главы 6.1):



Что нужно изменить в описании алгоритма, чтобы он продолжал работать при подобной нумерации? Каждая буква теперь идентифицируется не одним числом, а парой — номера слева и справа от нее. При этом чисел стало на одно больше, чем при прежнем способе нумерации.

Возьмем матрицу  $(n+1) \times (n+1) \times |N|$  и при инициализации будем заполнять не главную диагональ, а диагональ прямо над ней. Таким образом, мы начинаем наш алгоритм с определения значений  $M[i, j, A]$ , где  $j = i + 1$ . При этом наши дальнейшие действия в рамках алгоритма не изменятся.

Для примера 6.1.1 на шаге инициализации матрица выглядит следующим образом:

$$\begin{pmatrix} \emptyset & \{A\} & \emptyset & \emptyset & \emptyset & \emptyset & \emptyset \\ \emptyset & \emptyset & \{A\} & \emptyset & \emptyset & \emptyset & \emptyset \\ \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset & \emptyset \\ \emptyset & \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset \\ \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \{A\} & \emptyset \\ \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} \\ \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \emptyset \end{pmatrix}$$

А в результате работы алгоритма имеем:

$$\begin{pmatrix} \emptyset & \{A\} & \emptyset & \emptyset & \{S_1, S\} & \emptyset & \{S_1, S\} \\ \emptyset & \emptyset & \{A\} & \{S_1\} & \{S_2\} & \emptyset & \{S_2\} \\ \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \emptyset & \emptyset \\ \emptyset & \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} & \emptyset & \{S_2, S_3\} \\ \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \{A\} & \{S_1\} \\ \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \{B, S_2, S_3\} \\ \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \emptyset & \emptyset \end{pmatrix}$$

Мы представили входную строку в виде линейного графа, а на шаге инициализации получили его матрицу смежности. Добавление нового нетерминала в ячейку матрицы можно рассматривать как нахождение нового пути между соответствующими вершинами, выводимого из добавленного нетерминала. Таким образом, шаги алгоритма напоминают построение транзитивного замыкания графа. Различие заключается в том, что мы добавляем новые ребра только для тех пар нетерминалов, для которых существует соответствующее правило в грамматике.

Алгоритм можно обобщить и на произвольные графы с метками, рассматриваемые в этом курсе. При этом можно ослабить ограничение на форму входной грамматики: она должна находиться в ослабленной Нормальной Форме Хомского (5.3.3).

*Почему?*

Шаг инициализации в алгоритме теперь состоит из двух пунктов.

- Как и раньше, с помощью продукций вида

$$A \rightarrow a, \text{ где } A \in N, a \in \Sigma$$

заменяем терминалы на ребрах входного графа на множества нетерминалов, из которых они выводятся.

- Добавляем в каждую вершину петлю, помеченную множеством нетерминалов для которых в данной грамматике есть правила вида

$$A \rightarrow \varepsilon, \text{ где } A \in N.$$

Затем используем матрицу смежности получившегося графа (обозначим ее  $M$ ) в качестве начального значения. Дальнейший ход алгоритма можно описать псевдокодом, представленным в листинге 3.

---

**Algorithm 3** Алгоритм КС достижимости на основе CYK
 

---

```

1: function CONTEXTFREEPATHQUERYING( $G, \mathcal{G}$ )
2:    $n \leftarrow$  the number of nodes in  $\mathcal{G}$ 
3:    $M \leftarrow$  the modified adjacency matrix of  $\mathcal{G}$ 
4:    $P \leftarrow$  the set of production rules in  $G$ 
5:   while  $M$  is changing do
6:     for  $k \in 0..n$  do
7:       for  $i \in 0..n$  do
8:         for  $j \in 0..n$  do
9:           for all  $N_1 \in M[i, k], N_2 \in M[k, j]$  do
10:            if  $N_3 \rightarrow N_1 N_2 \in P$  then
11:               $M[i, j] += \{N_3\}$ 
12:   return  $M$ 

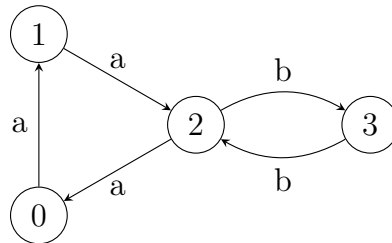
```

} перевернули

↖ U = ?

После завершения алгоритма, если в некоторой ячейке результирующей матрицы с номером  $(i, j)$  находятся стартовый нетерминал, то это означает, что существует путь из вершины  $i$  в вершину  $j$ , удовлетворяющий данной грамматике. Таким образом, полученная матрица является ответом для задачи достижимости для заданных графа и грамматики.

**Пример 6.2.1.** Рассмотрим работу алгоритма на графе

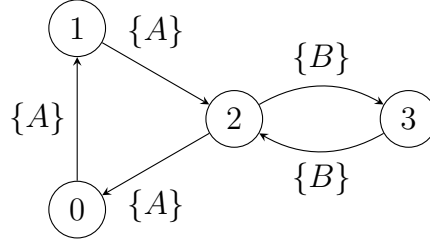


и грамматике:

$$\begin{array}{ll}
 S \rightarrow AB & A \rightarrow a \\
 S \rightarrow AS_1 & B \rightarrow b \\
 S_1 \rightarrow SB &
 \end{array}$$

Данный пример является классическим и еще не раз будет использоваться в рамках данного курса.

**Инициализация.** Заменяем терминалы на ребрах графа на нетерминалы, из которых они выводятся, и строим матрицу смежности получившегося графа:



$$\begin{pmatrix} \emptyset & \{A\} & \emptyset & \emptyset \\ \emptyset & \emptyset & \{A\} & \emptyset \\ \{A\} & \emptyset & \emptyset & \{B\} \\ \emptyset & \emptyset & \{B\} & \emptyset \end{pmatrix}$$

**Итерация 1.** Итерируемся по  $k$ ,  $i$  и  $j$ , пытаясь найти пары нетерминалов, для которых существуют правила вывода, их выводящие. Нам интересны следующие случаи:

- $k = 2, i = 1, j = 3$  :  $A \in M[1, 2], B \in M[2, 3]$ , так как в грамматике присутствует правило  $S \rightarrow AB$ , добавляем нетерминал  $S$  в ячейку  $M[1, 3]$ .
- $k = 3, i = 1, j = 2$  :  $S \in M[1, 3], B \in M[3, 2]$ , поскольку в грамматике есть правило  $S_1 \rightarrow SB$ , добавляем нетерминал  $S_1$  в ячейку  $M[1, 2]$ .

В остальных случаях либо какая-то из клеток пуста, либо не существует продукции в грамматике, выводящей данные два нетерминала.

Матрица после данной итерации:

$$\begin{pmatrix} \emptyset & \{A\} & \emptyset & \emptyset \\ \emptyset & \emptyset & \{A, S_1\} & \{S\} \\ \{A\} & \emptyset & \emptyset & \{B\} \\ \emptyset & \emptyset & \{B\} & \emptyset \end{pmatrix}$$

**Итерация 2.** Снова итерируемся по  $k$ ,  $i$ ,  $j$ . Рассмотрим случаи:



- $k = 1, i = 0, j = 2 : A \in M[0, 1], S_1 \in M[1, 2]$ , так как в грамматике присутствует правило  $S \rightarrow AS_1$ , добавляем нетерминал  $S$  в ячейку  $M[0, 2]$ .
- $k = 2, i = 0, j = 3 : S \in M[0, 2], B \in M[2, 3]$ , поскольку в грамматике есть правило  $S_1 \rightarrow SB$ , добавляем нетерминал  $S_1$  в ячейку  $M[0, 3]$ .

Матрица на данном шаге:

$$\begin{pmatrix} \emptyset & \{A\} & \{S\} & \{S_1\} \\ \emptyset & \emptyset & \{A, S_1\} & \{S\} \\ \{A\} & \emptyset & \emptyset & \{B\} \\ \emptyset & \emptyset & \{B\} & \emptyset \end{pmatrix}$$

**Итерация 3.** Рассматриваемые на данном шаге случаи:

- $k = 0, i = 2, j = 3 : A \in M[2, 0], S_1 \in M[0, 3]$ , так как в грамматике присутствует правило  $S \rightarrow AS_1$ , добавляем нетерминал  $S$  в ячейку  $M[2, 3]$ .
- $k = 3, i = 2, j = 2 : S \in M[2, 3], B \in M[3, 2]$ , поскольку в грамматике есть правило  $S_1 \rightarrow SB$ , добавляем нетерминал  $S_1$  в ячейку  $M[2, 2]$ .

Матрица после этой итерации:

$$\begin{pmatrix} \emptyset & \{A\} & \{S\} & \{S_1\} \\ \emptyset & \emptyset & \{A, S_1\} & \{S\} \\ \{A\} & \emptyset & \{S_1\} & \{B, S\} \\ \emptyset & \emptyset & \{B\} & \emptyset \end{pmatrix}$$

**Итерация 4.** Рассматриваемые случаи:

- $k = 2, i = 1, j = 2 : A \in M[1, 2], S_1 \in M[2, 2]$ , так как в грамматике присутствует правило  $S \rightarrow AS_1$ , добавляем нетерминал  $S$  в ячейку  $M[1, 2]$ .
- $k = 2, i = 1, j = 3 : S \in M[1, 2], B \in M[2, 3]$ , поскольку в грамматике есть правило  $S_1 \rightarrow SB$ , добавляем нетерминал  $S_1$  в ячейку  $M[1, 3]$ .

Матрица:

$$\begin{pmatrix} \emptyset & \{A\} & \{S\} & \{S_1\} \\ \emptyset & \emptyset & \{A, \mathbf{S}, S_1\} & \{S, \mathbf{S}_1\} \\ \{A\} & \emptyset & \{S_1\} & \{B, S\} \\ \emptyset & \emptyset & \{B\} & \emptyset \end{pmatrix}$$

**Итерация 5.** Рассмотрим на это шаге:

- $k = 1, i = 0, j = 3 : A \in M[0, 1], S_1 \in M[1, 3]$ , поскольку в грамматике есть правило  $S \rightarrow AS_1$ , добавляем нетерминал  $S$  в ячейку  $M[0, 3]$ .
- $k = 3, i = 0, j = 2 : S \in M[0, 3], B \in M[3, 2]$ , поскольку в грамматике есть правило  $S_1 \rightarrow SB$ , добавляем нетерминал  $S_1$  в ячейку  $M[0, 2]$ .

Матрица на этой итерации:

$$\begin{pmatrix} \emptyset & \{A\} & \{S, \mathbf{S}_1\} & \{S, S_1\} \\ \emptyset & \emptyset & \{A, S, S_1\} & \{S, S_1\} \\ \{A\} & \emptyset & \{S_1\} & \{B, S\} \\ \emptyset & \emptyset & \{B\} & \emptyset \end{pmatrix}$$

**Итерация 6.** Интересующие нас на этом шаге случаи:

- $k = 0, i = 2, j = 2 : A \in M[2, 0], S_1 \in M[0, 2]$ , поскольку в грамматике есть правило  $S \rightarrow AS_1$ , добавляем нетерминал  $S$  в ячейку  $M[2, 2]$ .
- $k = 2, i = 2, j = 3 : S \in M[2, 2], B \in M[2, 3]$ , поскольку в грамматике есть правило  $S_1 \rightarrow SB$ , добавляем нетерминал  $S_1$  в ячейку  $M[2, 3]$ .

Матрица после данного шага:

$$\begin{pmatrix} \emptyset & \{A\} & \{S, S_1\} & \{S, S_1\} \\ \emptyset & \emptyset & \{A, S, S_1\} & \{S, S_1\} \\ \{A\} & \emptyset & \{\mathbf{S}, S_1\} & \{B, S, \mathbf{S}_1\} \\ \emptyset & \emptyset & \{B\} & \emptyset \end{pmatrix}$$

На следующей итерации матрица не изменяется, поэтому заканчиваем работу алгоритма. В результате, если ячейка  $M[i, j]$  содержит стартовый нетерминал  $S$ , то существует путь из  $i$  в  $j$ , удовлетворяющий ограничениям, заданным грамматикой.

Можно заметить, что мы делаем много лишних итераций. Можно переписать алгоритм так, чтобы он не просматривал заведомо пустые ячейки. Данную модификацию предложил Й.Хеллингс в работе [31], также она реализована в работе [78].

Псевдокод алгоритма Хеллингса представлен в листинге 4.

---

**Algorithm 4** Алгоритм Хеллингса
 

---

```

1: function CONTEXTFREEPATHQUERYING( $G = \langle \Sigma, N, P, S \rangle$ ,  $\mathcal{G} = \langle V, E, L \rangle$ )
2:    $r \leftarrow \{(N_i, v, v) \mid v \in V \wedge N_i \rightarrow \varepsilon \in P\} \cup \{(N_i, v, u) \mid (v, t, u) \in E \wedge N_i \rightarrow t \in P\}$ 
3:    $m \leftarrow r$ 
4:   while  $m \neq \emptyset$  do
5:      $(N_i, v, u) \leftarrow m.\text{pick}()$ 
6:     for  $(N_j, v', v) \in r$  do
7:       for  $N_k \rightarrow N_j N_i \in P$  таких что  $((N_k, v', u) \notin r)$  do
8:          $m \leftarrow m \cup \{(N_k, v', u)\}$ 
9:          $r \leftarrow r \cup \{(N_k, v', u)\}$ 
10:    for  $(N_j, u, v') \in r$  do
11:      for  $N_k \rightarrow N_i N_j \in P$  таких что  $((N_k, v, v') \notin r)$  do
12:         $m \leftarrow m \cup \{(N_k, v, v')\}$ 
13:         $r \leftarrow r \cup \{(N_k, v, v')\}$ 
14:   return  $r$ 

```

---

**Пример 6.2.2.** Запустим алгоритм Хеллингса на нашем примере.

**Инициализация**

$$m = r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2)\}$$

**Итерации внешнего цикла.** Будем считать, что  $r$  и  $m$  — упорядоченные списки и  $\text{pick}$  возвращает его голову, оставляя хвост. Новые элементы добавляются в конец.

*pop? first? head?*

1. Обработываем  $(A, 0, 1)$ . Ни один из вложенных циклов не найдёт новых путей, так как для рассматриваемого ребра есть только две возможности достроить путь:  $2 \xrightarrow{A} 0 \xrightarrow{A} 1$  и  $0 \xrightarrow{A} 1 \xrightarrow{A} 2$  и ни одна из соответствующих строк не выводится в заданной граммтике.

2. Перед началом итерации

$$m = \{(A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2)\},$$

$r$  не изменилось. Обрабатываем  $(A, 1, 2)$ . В данной ситуации второй цикл найдёт тройку  $(B, 2, 3)$  и соответствующее правило  $S \rightarrow A B$ . Это значит, что и в  $m$  и в  $r$  добавится тройка  $(S, 1, 3)$ .

3. Перед началом итерации

$$m = \{(A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3)\}.$$

Обрабатываем  $(A, 2, 0)$ . Внутринние циклы ничего не найдут, новых путей не появится.

4. Перед началом итерации

$$m = \{(B, 2, 3), (B, 3, 2), (S, 1, 3)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3)\}.$$

Обрабатываем  $(B, 2, 3)$ . Первый цикл мог бы найти  $(A, 1, 2)$ , однако при проверке во вложенном цикле выяснится, что  $(S, 1, 3)$  уже найдена. В итоге, на данной итерации новых путей не появится.

5. Перед началом итерации

$$m = \{(B, 3, 2), (S, 1, 3)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3)\}.$$

Обрабатываем  $(B, 3, 2)$ . Первый цикл найдёт  $(S, 1, 3)$  и соответствующее правило  $S_1 \rightarrow S B$ . Это значит, что и в  $m$  и в  $r$  добавится тройка  $(S_1, 1, 2)$ .

6. Перед началом итерации

$$m = \{(S, 1, 3), (S_1, 1, 2)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3), (S_1, 1, 2)\}.$$

Обрабатываем  $(S, 1, 3)$ . Второй цикл мог бы найти  $(B, 3, 2)$ , однако при проверке во вложенном цикле выяснится, что  $(S_1, 1, 2)$  уже найдена. В итоге, на данной итерации новых путей не появится.

7. Перед началом итерации

$$m = \{(S_1, 1, 2)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3), (S_1, 1, 2)\}.$$

Обрабатываем  $(S_1, 1, 2)$ . Первый цикл найдёт  $(A, 0, 1)$  и соответствующее правило  $S \rightarrow A S_1$ . Это значит, что и в  $m$  и в  $r$  добавится тройка  $(S, 0, 2)$ .

8. Перед началом итерации

$$m = \{(S, 0, 2)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3), (S_1, 1, 2), (S, 0, 2)\}.$$

Обрабатываем  $(S, 0, 2)$ . Найдено:  $(B, 2, 3)$  и соответствующее правило  $S_1 \rightarrow S B$ . В  $m$  и в  $r$  добавится тройка  $(S_1, 0, 3)$ .

9. Перед началом итерации

$$m = \{(S_1, 0, 3)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3), (S_1, 1, 2), (S, 0, 2), (S_1, 0, 3)\}.$$

Обрабатываем  $(S_1, 0, 3)$ . Найдено:  $(A, 2, 0)$  и соответствующее правило  $S \rightarrow A S_1$ . В  $m$  и в  $r$  добавится тройка  $(S, 2, 3)$ .

10. Перед началом итерации

$$m = \{(S, 2, 3)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3), (S_1, 1, 2), (S, 0, 2), (S_1, 0, 3), (S, 2, 3)\}.$$

Обрабатываем  $(S, 2, 3)$ . Найдено:  $(B, 3, 2)$  и соответствующее правило  $S_1 \rightarrow S B$ . В  $m$  и в  $r$  добавится тройка  $(S_1, 2, 2)$ .

11. Перед началом итерации

$$m = \{(S_1, 2, 2)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3), (S_1, 1, 2), (S, 0, 2), (S_1, 0, 3), (S, 2, 3), (S_1, 2, 2)\}.$$

Обрабатываем  $(S_1, 2, 2)$ . Найдено:  $(A, 1, 2)$  и соответствующее правило  $S \rightarrow A S_1$ . В  $m$  и в  $r$  добавится тройка  $(S, 1, 2)$ .

12. Перед началом итерации

$$m = \{(S, 1, 2)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3), (S_1, 1, 2), (S, 0, 2), (S_1, 0, 3), (S, 2, 3), (S_1, 2, 2), (S, 1, 2)\}.$$

Обрабатываем  $(S, 1, 2)$ . Найдено:  $(B, 2, 3)$  и соответствующее правило  $S_1 \rightarrow S B$ . В  $m$  и в  $r$  добавится тройка  $(S_1, 1, 3)$ .

13. Перед началом итерации

$$m = \{(S_1, 1, 3)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3), (S_1, 1, 2), (S, 0, 2), (S_1, 0, 3), (S, 2, 3), (S_1, 2, 2), (S, 1, 2), (S_1, 1, 3)\}.$$

Обрабатываем  $(S_1, 1, 3)$ . Найдено:  $(A, 0, 1)$  и соответствующее правило  $S \rightarrow A S_1$ . В  $m$  и в  $r$  добавится тройка  $(S, 0, 3)$ .

14. Перед началом итерации

$$m = \{(S, 0, 3)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3), (S_1, 1, 2), (S, 0, 2), (S_1, 0, 3), (S, 2, 3), (S_1, 2, 2), (S, 1, 2), (S_1, 1, 3), (S, 0, 3)\}.$$

Обрабатываем  $(S, 0, 3)$ . Найдено:  $(B, 3, 2)$  и соответствующее правило  $S_1 \rightarrow S B$ . В  $m$  и в  $r$  добавится тройка  $(S_1, 0, 2)$ .

15. Перед началом итерации

$$m = \{(S_1, 0, 2)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3), (S_1, 1, 2), (S, 0, 2), (S_1, 0, 3), (S, 2, 3), (S_1, 2, 2), (S, 1, 2), (S_1, 1, 3), (S, 0, 3), (S_1, 0, 2)\}.$$

Обрабатываем  $(S_1, 0, 2)$ . Найдено:  $(A, 2, 0)$  и соответствующее правило  $S \rightarrow A S_1$ . В  $m$  и в  $r$  добавится тройка  $(S, 2, 2)$ .

16. Перед началом итерации

$$m = \{(S, 2, 2)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3), (S_1, 1, 2), (S, 0, 2), \\ (S_1, 0, 3), (S, 2, 3), (S_1, 2, 2), (S, 1, 2), (S_1, 1, 3), (S, 0, 3), (S_1, 0, 2), \\ (S, 2, 2)\}.$$

Обрабатываем  $(S, 2, 2)$ . Найдено:  $(B, 2, 3)$  и соответствующее правило  $S_1 \rightarrow S B$ . В  $m$  и в  $r$  добавится тройка  $(S_1, 2, 3)$ .

17. Перед началом итерации

$$m = \{(S_1, 2, 3)\},$$

$$r = \{(A, 0, 1), (A, 1, 2), (A, 2, 0), (B, 2, 3), (B, 3, 2), (S, 1, 3), (S_1, 1, 2), (S, 0, 2), \\ (S_1, 0, 3), (S, 2, 3), (S_1, 2, 2), (S, 1, 2), (S_1, 1, 3), (S, 0, 3), (S_1, 0, 2), \\ (S, 2, 2), (S_1, 2, 3)\}.$$

Обрабатываем  $(S_1, 2, 3)$ . Могло бы быть найдено:  $(A, 1, 2)$  и соответствующее правило  $S \rightarrow A S_1$ , однако тройка  $(S, 1, 3)$  уже есть в  $r$ . А значит никаких новых троек найдено не будет и  $m$  становится пустым. Это была последняя итерация внешнего цикла, в  $r$  на текущий момент уже содержится всё ршение.

Как можно заметить, количество итераций внешнего цикла также получилось достаточно большим. Проверьте, зависит ли оно от порядка обработки элементов из  $m$ . При этом внутренние циклы в нашем случае достаточно короткие, так как просматриваются только “существенные” элементы и избегается дублирование.





## Глава 7

# КС и конъюнктивная достижимость через произведение матриц

В данном разделе мы рассмотрим алгоритм решения задачи контекстно-свободной и конъюнктивной достижимости, основанный на произведении матриц. Будет показано, что при использовании конъюнктивных грамматик, представленный алгоритм находит переаппроксимацию истинного решения задачи.

### 7.1 КС достижимость через произведение матриц

В главе 6.2 был изложен алгоритм для решения задачи КС достижимости на основе СΥК. Заметим, что обход матрицы напоминает умножение матриц в ячейках которых множества нетерминалов:

$$M_3 = M_1 \times M_2$$
$$M_3[i, j] = \sum_{k=1}^n M[i, k] * M[k, j]$$

, где сумма — это объединение множеств:

$$\sum_{k=1}^n = \bigcup_{k=1}^n$$

, а поэлементное умножение определено следующим образом:

$$S_1 * S_2 = \{N_1^0 \dots N_1^m\} * \{N_2^0 \dots N_2^l\} = \{N_3 \mid (N_3 \rightarrow N_1^i N_2^j) \in P\}.$$

Таким образом, алгоритм решения задачи КС достижимости может быть выражена в терминах перемножения матриц над полукольцом с соответствующими операциями.

Для частного случая этой задачи, синтаксического анализа линейного входа, существует алгоритм Валианта [68], использующий эту идею. Однако он не обобщается на графы из-за того, что существенно использует возможность упорядочить обход матрицы (см. разницу в СҮК для линейного случая и для графа). Поэтому, хотя для линейного случая алгоритм Валианта является алгоритмом синтаксического анализа для произвольных КС грамматик за субкубическое время, его обобщение до задачи КС достижимости в произвольных графах с сохранением асимптотики является нетривиальной задачей [75]. В настоящее время алгоритм с субкубической сложностью получен только для частного случая — языка Дика с одним типом скобок — Филипом Брэдфордом [16].

В случае с линейным входом, отдельного внимания заслуживает работа Лиллиан Ли (Lillian Lee) [42], где она показывает, что задача перемножения матриц сводима к синтаксическому анализу линейного входа. Аналогичных результатов для графов на текущий момент не известно.

Поэтому рассмотрим более простую идею, изложенную в статье Рустама Азимова [6]: будем строить транзитивное замыкание графа через наивное (не через возведение в квадрат) умножение матриц.

Пусть  $\mathcal{G} = (V, E)$  — входной граф и  $G = (N, \Sigma, P)$  — входная грамматика. Тогда алгоритм может быть сформулирован как представлено в листинге 5.

**Algorithm 5** Context-free recognizer for graphs

---

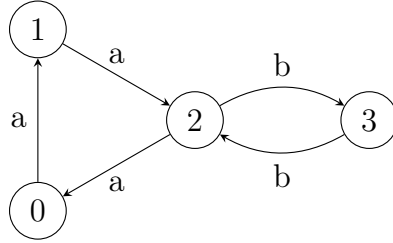
```

1: function CONTEXTFREEPATHQUERYING( $\mathcal{G}$ ,  $G$ )
2:    $n \leftarrow$  количество узлов в  $\mathcal{G}$ 
3:    $E \leftarrow$  направленные ребра в  $\mathcal{G}$ 
4:    $P \leftarrow$  набор продукций из  $G$ 
5:    $T \leftarrow$  матрица  $n \times n$ , в которой каждый элемент  $\emptyset$ 
6:   for all  $(i, x, j) \in E$  do ▷ Инициализация матрицы
7:      $T_{i,j} \leftarrow T_{i,j} \cup \{A \mid (A \rightarrow x) \in P\}$ 
8:   for all  $i \in 0 \dots n - 1$  do ▷ Добавление петель для нетерминалов,
порождающих пустую строку
9:      $T_{i,i} \leftarrow T_{i,i} \cup \{A \in N \mid A \rightarrow \varepsilon\}$ 
10:  while матрица  $T$  меняется do
11:     $T \leftarrow T \cup (T \times T)$  ▷ Вычисление транзитивного замыкания
12:  return  $T$ 

```

---

**Пример 7.1.1** (Пример работы). Пусть есть граф  $\mathcal{G}$ :



и грамматика  $G$ :

$$\begin{array}{ll}
 S \rightarrow AB & A \rightarrow a \\
 S \rightarrow AS_1 & B \rightarrow b \\
 S_1 \rightarrow SB
 \end{array}$$

Пусть  $T_i$  — матрица, полученная из  $T$  после применения цикла, описанного в строках **8-9** алгоритма 5,  $i$  раз. Тогда  $T_0$ , полученная напрямую из графа, выглядит следующим образом:

$$T_0 = \begin{pmatrix} \emptyset & \{A\} & \emptyset & \{B\} \\ \emptyset & \emptyset & \{A\} & \emptyset \\ \{A\} & \emptyset & \emptyset & \emptyset \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix}$$

Далее показано получение матрицы  $T_1$ .

$$T_0 \times T_0 = \begin{pmatrix} \emptyset & \emptyset & \emptyset & \emptyset \\ \emptyset & \emptyset & \emptyset & \emptyset \\ \emptyset & \emptyset & \emptyset & \{S\} \\ \emptyset & \emptyset & \emptyset & \emptyset \end{pmatrix}$$

$$T_1 = T_0 \cup (T_0 \times T_0) = \begin{pmatrix} \emptyset & \{A\} & \emptyset & \{B\} \\ \emptyset & \emptyset & \{A\} & \emptyset \\ \{A\} & \emptyset & \emptyset & \{S\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix}$$

После первой итерации цикла нетерминал в ячейку  $T[2, 3]$  добавился нетерминал  $S$ . Это означает, что существует такой путь  $\pi$  из вершины 2 в вершину 3 в графе  $\mathcal{G}$ , что  $S \xrightarrow{*} \omega(\pi)$ . В данном примере путь состоит из двух ребер  $2 \xrightarrow{a} 0$  и  $0 \xrightarrow{b} 3$ , так что  $S \xrightarrow{*} ab$ .

Вычисление транзитивного замыкания заканчивается через  $k$  итераций, когда достигается неподвижная точка процесса:  $T_{k-1} = T_k$ . Для данного примера  $k = 13$ , так как  $T_{13} = T_{12}$ . Весь процесс работы алгоритма (все матрицы  $T_i$ ) показан ниже (на каждой итерации новые элементы выделены жирным).

$$\begin{aligned}
T_2 &= \begin{pmatrix} \emptyset & \{A\} & \emptyset & \{B\} \\ \emptyset & \emptyset & \{A\} & \emptyset \\ \{A, S_1\} & \emptyset & \emptyset & \{S\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix} & T_3 &= \begin{pmatrix} \emptyset & \{A\} & \emptyset & \{B\} \\ \{S\} & \emptyset & \{A\} & \emptyset \\ \{A, S_1\} & \emptyset & \emptyset & \{S\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix} \\
T_4 &= \begin{pmatrix} \emptyset & \{A\} & \emptyset & \{B\} \\ \{S\} & \emptyset & \{A\} & \{S_1\} \\ \{A, S_1\} & \emptyset & \emptyset & \{S\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix} & T_5 &= \begin{pmatrix} \emptyset & \{A\} & \emptyset & \{B, S\} \\ \{S\} & \emptyset & \{A\} & \{S_1\} \\ \{A, S_1\} & \emptyset & \emptyset & \{S\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix} \\
T_6 &= \begin{pmatrix} \{S_1\} & \{A\} & \emptyset & \{B, S\} \\ \{S\} & \emptyset & \{A\} & \{S_1\} \\ \{A, S_1\} & \emptyset & \emptyset & \{S\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix} & T_7 &= \begin{pmatrix} \{S_1\} & \{A\} & \emptyset & \{B, S\} \\ \{S\} & \emptyset & \{A\} & \{S_1\} \\ \{A, S_1, S\} & \emptyset & \emptyset & \{S\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix} \\
T_8 &= \begin{pmatrix} \{S_1\} & \{A\} & \emptyset & \{B, S\} \\ \{S\} & \emptyset & \{A\} & \{S_1\} \\ \{A, S_1, S\} & \emptyset & \emptyset & \{S, S_1\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix} & T_9 &= \begin{pmatrix} \{S_1\} & \{A\} & \emptyset & \{B, S\} \\ \{S\} & \emptyset & \{A\} & \{S_1, S\} \\ \{A, S_1, S\} & \emptyset & \emptyset & \{S, S_1\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix} \\
T_{10} &= \begin{pmatrix} \{S_1\} & \{A\} & \emptyset & \{B, S\} \\ \{S, S_1\} & \emptyset & \{A\} & \{S_1, S\} \\ \{A, S_1, S\} & \emptyset & \emptyset & \{S, S_1\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix} & T_{11} &= \begin{pmatrix} \{S_1, S\} & \{A\} & \emptyset & \{B, S\} \\ \{S, S_1\} & \emptyset & \{A\} & \{S_1, S\} \\ \{A, S_1, S\} & \emptyset & \emptyset & \{S, S_1\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix} \\
T_{12} &= \begin{pmatrix} \{S_1, S\} & \{A\} & \emptyset & \{B, S, S_1\} \\ \{S, S_1\} & \emptyset & \{A\} & \{S_1, S\} \\ \{A, S_1, S\} & \emptyset & \emptyset & \{S, S_1\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix} & T_{13} &= \begin{pmatrix} \{S_1, S\} & \{A\} & \emptyset & \{B, S, S_1\} \\ \{S, S_1\} & \emptyset & \{A\} & \{S_1, S\} \\ \{A, S_1, S\} & \emptyset & \emptyset & \{S, S_1\} \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix}
\end{aligned}$$

Таким образом, результат алгоритма 5 для нашего примера — это матрица  $T_{13} = T_{12}$ . Заметим, что для данного алгоритма приведённый пример также является худшим случаем: на каждой итерации в матрицу добавляется ровно один нетерминал, при том, что необходимо заполнить порядка  $O(n^2)$  ячеек.

### 7.1.1 Расширение алгоритма для конъюнктивных грамматик

Матричный алгоритм для конъюнктивных грамматик отличается от алгоритма 5 для контекстно-свободных грамматик только операцией умножения матриц, в остальном алгоритм остается без изменений. Определим операцию умножения матриц.

**Определение 7.1.1.** Пусть  $M_1$  и  $M_2$  матрицы размера  $n$ . Определим операцию  $\circ$  следующим образом:

$$M_1 \circ M_2 = M_3,$$

$$M_3[i, j] = \{A \mid \exists (A \rightarrow B_1 C_1 \& \dots \& B_m C_m) \in P, (B_k, C_k) \in d[i, j] \forall k = 1, \dots, m\}$$

, где

$$d[i, j] = \bigcup_{k=1}^n M_1[i, k] \times M_2[k, j].$$

Важно заметить, что алгоритм для конъюнктивных грамматик, в отличие от алгоритма для контекстно-свободных грамматик, дает лишь верхнюю оценку ответа. То есть все нетерминалы, которые должны быть в ячейках матрицы результата, содержатся там, но вместе с ними содержатся и лишние нетерминалы. Рассмотрим пример, иллюстрирующий появление лишних нетерминалов.

**Пример 7.1.2.** Грамматика  $G$ :

$$S \rightarrow AB\&DC$$

$$C \rightarrow c$$

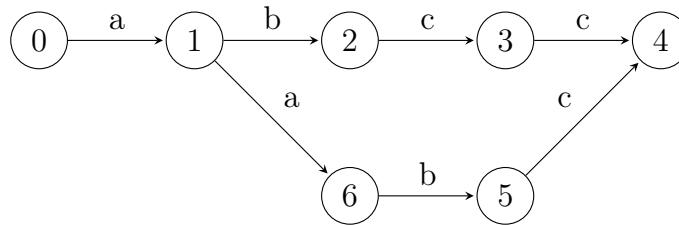
$$A \rightarrow a$$

$$D \rightarrow DC \mid b$$

$$B \rightarrow BC \mid b$$

Очевидно, что грамматика  $G$  задает язык из одного слова  $L(G) = \{abc\} = \{abc^*\} \cap \{a^*bc\}$ .

Пусть есть граф  $\mathcal{G}$ :



Применяя алгоритм, получим, что существует путь из вершины 0 в вершину 4, выводимый из нетерминала  $S$ . Однако очевидно, что в графе такого пути нет. Такое поведение алгоритма наблюдается из-за того, что существует путь “abcc”, соответствующий  $L(AB) = \{abc^*\}$  и путь “aabc”, соответствующий  $L(DC) = \{a^*bc\}$ , но они различны. Однако алгоритм не может это проверить, так как оперирует понятием достижимости между вершинами, а не наличием различных путей. Более того, в общем случае для конъюнктивных грамматик такую проверку реализовать нельзя. Поэтому для классической семантики достижимости с ограничениями в терминах конъюнктивных грамматик результат работы алгоритма является оценкой сверху.

Существует альтернативная семантика, когда мы трактуем конъюнкцию в правой части правил как конъюнкцию в Datalog (подробнее о Datalog в параграфе ??). Т.е. если есть правило  $S \rightarrow AB \& DC$ , то должен быть путь принадлежащий языку  $L(AB)$  и путь принадлежащий языку  $L(DC)$ . В такой семантике алгоритм дает точный ответ.

Подробнее алгоритм описан в статье Рустама Азимова и Семёна Григорьева [7]. Стоит также отметить, что обобщения данного алгоритма для булевых грамматик не существует, хотя и существует частное решение для случая, когда граф не содержит циклов (является DAG-ом), предложенное Екатериной Шеметовой [61].

## 7.2 Особенности реализации

Алгоритмы, описанные выше, удобны с точки зрения реализации тем, что могут быть эффективно реализованы с использованием высокопроизводительных библиотек линейной алгебры, которые эксплуатируют возможности параллельных вычислений на современных CPU и GPU [44]. Это позволяет с минимальными затратами получить эффективную параллельную реализацию алгоритма для решения задачи КС достижимости в графах. Благодаря этому, хотя асимптотически приведенные алгоритмы имеют большую сложность чем, скажем, алгоритм Хеллингса, в результате эффективного распараллеливания на практике они работают быстрее однопоточных алгоритмов с лучшей сложностью.

Далее рассмотрим некоторые детали, упрощающие реализацию с использованием современных библиотек и аппаратного обеспечения.

Так как множество нетерминалов и правил конечно, то мы можем свести представленный выше алгоритм к булевым матрицам: для каждого нетерминала заведём матрицу, такую что в ячейке стоит 1 тогда и только тогда, когда в исходной матрице в соответствующей ячейке содержится этот нетерминал. Тогда перемножение пары таких матриц, соответствующих нетерминалам  $A$  и  $B$ , соответствует построению путей, выводимых из нетерминалов, для которых есть правила с правой частью вида  $AB$ .

**Пример 7.2.1.** Представим в виде набора булевых матриц следующую матрицу:

$$T_0 = \begin{pmatrix} \emptyset & \{A\} & \emptyset & \{B\} \\ \emptyset & \emptyset & \{A\} & \emptyset \\ \{A\} & \emptyset & \emptyset & \emptyset \\ \{B\} & \emptyset & \emptyset & \emptyset \end{pmatrix}$$

Тогда:

$$T_{0\_A} = \begin{pmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \quad T_{0\_B} = \begin{pmatrix} 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 \end{pmatrix}$$

Тогда при наличии правила  $S \rightarrow AB$  в грамматике получим:

$$T_{1\_S} = T_{0\_A} \times T_{0\_B} = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \end{pmatrix}$$

Алгоритм же может быть переформулирован так, как показано в листинге 6. Такой взгляд на алгоритм позволяет использовать для его реализации существующие высокопроизводительные библиотеки для работы с булевыми матрицами (например M4RI<sup>1</sup> [2]) или библиотеки для линейной алгебры (например CUSP [21]).

---

**Listing 6** Context-free path querying algorithm. Boolean matrix version

---

```

1: function EVALCFPQ( $D = (V, E), G = (N, \Sigma, P)$ )
2:    $n \leftarrow |V|$ 
3:    $T \leftarrow \{T^{A_i} \mid A_i \in N, T^{A_i} \text{ is a matrix } n \times n, T_{k,l}^{A_i} \leftarrow \text{false}\}$ 
4:   for all  $(i, x, j) \in E, A_k \mid A_k \rightarrow x \in P$  do  $T_{i,j}^{A_k} \leftarrow \text{true}$ 
5:   for  $A_k \mid A_k \rightarrow \varepsilon \in P$  do  $T_{i,i}^{A_k} \leftarrow \text{true}$ 
6:   while any matrix in  $T$  is changing do
7:     for  $A_i \rightarrow A_j A_k \in P$  do  $T^{A_i} \leftarrow T^{A_i} + (T^{A_j} \times T^{A_k})$ 
8:   return  $T$ 

```

---

С другой стороны, для запросов, выражимых в терминах граммтик с небольшим количеством нетерминалов, практически может быть выгодно представлять множества нетерминалов в ячейке матрицы в виде битового вектора следующим образом. Нумеруем все нетерминалы с нуля, в векторе стоит 1 на позиции  $i$ , если в множестве есть нетерминал с номером  $i$ . Таким образом, в каждой ячейке хранится битовый вектор длины  $|N|$ . Тогда операция умножения определяется следующим образом:

$$v_1 \times v_2 = \{v \mid \exists (v, v_3) \in P, \text{append}(v_1, v_2) \& v_3 = v_3\},$$

---

<sup>1</sup>M4RI — одна из высокопроизводительных библиотек для работы с логическими матрицами на CPU. Реализует Метод Четырёх Русских. Исходный код библиотеки: <https://bitbucket.org/malb/m4ri/src/master/>. Дата посещения: 30.03.2020.



где  $\&$  — побитовое “и”.

Правила надо кодировать соответственно: продукция это пара, где первый элемент — битовый вектор длины  $|N|$  с единственной единицей в позиции, соответствующей нетерминалу в правой части, а второй элемент — вектор длины  $2|N|$ , с двумя единицами кодирующими первый и второй нетерминалы, соответственно.

**Пример 7.2.2.** Пусть  $N = \{S, A, B\}$  и в грамматике есть продукция  $S \rightarrow AB$ . Тогда занумеруем нетерминалы  $(S, 0), (A, 1), (B, 2)$ . Продукция примет вид  $[1, 0, 0] \rightarrow [0, 1, 0, 0, 0, 1]$ . Матрица  $T_0$  примет вид (здесь “.” означает, что в ячейке стоит  $[0, 0, 0]$ ):

$$T_0 = \begin{pmatrix} . & [0, 1, 0] & . & [0, 0, 1] \\ . & . & [0, 1, 0] & . \\ [0, 1, 0] & . & . & . \\ [0, 0, 1] & . & . & . \end{pmatrix}$$

После выполнения умножения получим:

$$T_1 = T_0 + T_0 \times T_0 = \begin{pmatrix} . & [0, 1, 0] & . & [0, 0, 1] \\ . & . & [0, 1, 0] & . \\ [0, 1, 0] & . & . & [1, 0, 0] \\ [0, 0, 1] & . & . & . \end{pmatrix}$$

На практике в роли векторов могут выступать беззнаковые целые числа. Например, 32 бита под ячейки в матрице и 64 бита под правила (или 8 и 16, если позволяет количество нетерминалов в грамматике, или 16 и 32). Тогда умножение выражается через битовые операции и сравнение, что довольно эффективно с точки зрения вычислений.

Для небольших запросов такой подход к реализации может оказаться быстрее — в данном случае скорость зависит от деталей. Минус подхода в том, что нет возможности использовать готовые библиотеку линейной алгебры без предварительной модификации. Только если они не являются параметризуемыми и не позволяют задать собственный тип и собственную операцию умножения и сложения (иными словами, собственное полукольцо). Такую возможность предусматривает, например, стандарт GraphBLAS<sup>2</sup> и, соответственно,

<sup>2</sup>GraphBLAS — открытый стандарт, описывающий набор примитивов и операций, необходимый для реализации графовых алгоритмов в терминах линейной алгебры. Web-страница проекта: <https://github.com/gunrock/graphblast>. Дата доступа: 30.03.2020.

его реализации, такие как SuiteSparse<sup>3</sup> [23].

Также стоит заметить, что при работе с реальными графами матрицы как правило оказываются разреженными, а значит необходимо использовать соответствующие представления матриц (CRS, покоординатное, Quad Tree [26]) и библиотеки, работающие с таким представлениями.

Однако даже при использовании разреженных матриц, могут возникнуть проблемы с размером реальных данных и объёмом памяти. Например, для вычислений на GPGPU лучше всего, когда все нужные для вычисления матрицы помещаются на одну карту. Тогда можно свести обмен данными между хостом и графическим сопроцессором к минимуму. Если не помещаются все, то нужно, чтобы помещалась хотя бы тройка непосредственно обрабатываемых матриц (два операнда и результат). В самом тяжёлом случае в памяти не удаётся разместить даже операнды целиком и тогда приходится прибегать к поблочному умножению матриц.

Отдельной инженерной проблемой является масштабирование алгоритмов на несколько вычислительных узлов, как на несколько CPU, так и на несколько GPGPU.

Важным свойством рассмотренного алгоритма является то, что описанные проблемы с объёмом памяти и масштабированием могут эффективно решаться в рамках библиотек линейной алгебры общего назначения, что избавляет от необходимости создавать специализированные решения для конкретных задач.

---

<sup>3</sup>SuiteSparse — это специализированное программное обеспечение для работы с разреженными матрицами, которое включает в себя реализацию GraphBLAS API. Web-страница проекта: <http://faculty.cse.tamu.edu/davis/suitesparse.html>. Дата доступа: 30.03.2020.

## Глава 8

# КС достижимость через тензорное произведение

Предыдущий подход позволяет выразить задачу поиска путей с ограничениями в терминах формальных языков через набор матричных операций. Это позволяет использовать высокопроизводительные библиотеки, массовопараллельные архитектуры и другие готовые решения для линейной алгебры. Однако, такой подход требует, чтобы грамматика находилась в ослабленной нормальной форме Хомского, что приводит к её разрастанию. Можно ли как-то избежать этого?

В данном разделе мы предложим альтернативное сведение задачи поиска путей к матричным операциям. В результате мы сможем избежать преобразования грамматики в ОНФХ, однако, матрицы, с которыми нам придётся работать, будут существенно большего размера.

В основе подхода лежит использование рекурсивных сетей или рекурсивных автоматов в качестве представления контекстно-свободных грамматик и использование тензорного (прямого) произведения для нахождения пересечения автоматов.

### 8.1 Рекурсивные автоматы и сети

Рекурсивный автомат или сеть — это представление контекстно-свободных грамматик, обобщающее конечные автоматы. В нашей работе мы будем придерживаться термина **рекурсивный автомат**. Классическое определение рекурсивного автомата выглядит следующим образом.

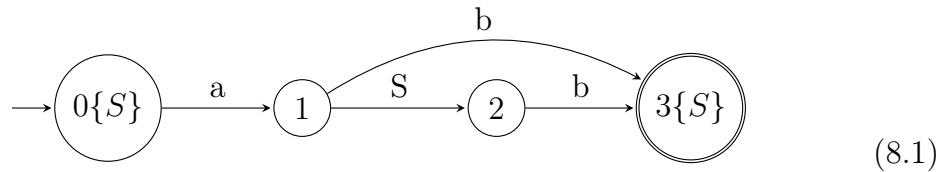
**Определение 8.1.1.** Рекурсивный автомат — это кортеж вида  $\langle N, \Sigma, S, D \rangle$ , где

- $N$  — нетерминальный алфавит;
- $\Sigma$  — терминальный алфавит;
- $S$  — стартовый нетерминал;
- $D$  — конечный автомат над  $N \cup \Sigma$  в котором стартовые и финальные состояния помечены подмножествами  $N$ .

Построим рекурсивный автомат для грамматики  $G$ :

$$S \rightarrow aSb$$

$$S \rightarrow ab$$



Используем стандартные обозначения для стартовых и финальных состояний. Дополнительно в стартовых и финальных состояниях укажем нетерминалы, для которых эти состояния стартовые/финальные.

В некоторых случаях рекурсивный автомат можно рассматривать как конечный автомат над смешанным алфавитом. Именно такой взгляд мы будем использовать при изложении алгоритма.

## 8.2 Тензорное произведение

Теперь перейдём к графам. Сперва дадим классическое определение тензорного произведения двух неориентированных графов.

**Определение 8.2.1.** Пусть даны два графа:  $\mathcal{G}_1 = \langle V_1, E_1 \rangle$  и  $\mathcal{G}_2 = \langle V_2, E_2 \rangle$ . Тензорным произведением этих графов будем называть граф  $\mathcal{G}_3 = \langle V_3, E_3 \rangle$ , где  $V_3 = V_1 \times V_2$ ,  $E_3 = \{((v_1, v_2), (u_1, u_2)) \mid (v_1, u_1) \in E_1 \text{ и } (v_2, u_2) \in E_2\}$ .

Иными словами, тензорным произведением двух графов является граф, такой что:

1. множество вершин — это прямое произведение множеств вершин исходных графов;
2. ребро между вершинами  $v = (v_1, v_2)$  и  $u = (u_1, u_2)$  существует тогда и только тогда, когда существуют рёбра между парами вершин  $v_1, u_1$  и  $v_2, u_2$  в соответствующих графах.

Для того, чтобы построить тензорное произведение ориентированных графов, необходимо в предыдущем определении, в условии существования ребра в результирующем графе, дополнительно потребовать, чтобы направления рёбер совпадали. Данное требование получается естественным образом, если считать, что пары вершин, задающие ребро, упорядочены, поэтому формальное определение отличаться не будет.

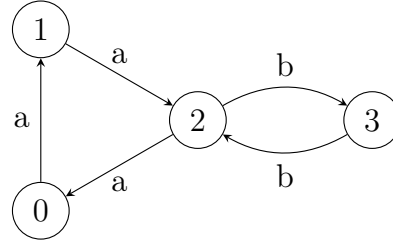
Осталось добавить метки к рёбрам. Это приведёт к логичному усилению требования к существованию ребра: метки рёбер в исходных графах должны совпадать. Таким образом, мы получаем следующее определение тензорного произведения ориентированных графов с метками на рёбрах.

**Определение 8.2.2.** Пусть даны два ориентированных графа с метками на рёбрах:  $\mathcal{G}_1 = \langle V_1, E_1, L_1 \rangle$  и  $\mathcal{G}_2 = \langle V_2, E_2, L_2 \rangle$ . Тензорным произведением этих графов будем называть граф  $\mathcal{G}_3 = \langle V_3, E_3, L_3 \rangle$ , где  $V_3 = V_1 \times V_2$ ,  $E_3 = \{((v_1, v_2), l, (u_1, u_2)) \mid (v_1, l, u_1) \in E_1 \text{ и } (v_2, l, u_2) \in E_2\}$ ,  $L_3 = L_1 \cap L_2$ .

Нетрудно заметить, что матрица смежности графа  $\mathcal{G}_3$  равна тензорному произведению матриц смежностей исходных графов  $\mathcal{G}_1$  и  $\mathcal{G}_2$ .

**Пример 8.2.1.** Рассмотрим пример. В качестве одного из графов возьмём рекурсивный автомат, построенный ранее (изображение 8.1). Его матрица смежности выглядит следующим образом.

$$M_1 = \begin{pmatrix} \cdot & [a] & \cdot & \cdot \\ \cdot & \cdot & [S] & [b] \\ \cdot & \cdot & \cdot & [b] \\ \cdot & \cdot & \cdot & \cdot \end{pmatrix}$$



(8.2)

Второй граф представлен на изображении 8.2. Его матрица смежности имеет следующий вид.

$$M_2 = \begin{pmatrix} \cdot & [a] & \cdot & \cdot \\ \cdot & \cdot & [a] & \cdot \\ [a] & \cdot & \cdot & [b] \\ \cdot & \cdot & [b] & \cdot \end{pmatrix}$$

Теперь вычислим  $M_1 \otimes M_2$ .

$$M_3 = M_1 \otimes M_2 = \begin{pmatrix} \cdot & [a] & \cdot & [b] \\ \cdot & \cdot & [S] & [b] \\ \cdot & \cdot & \cdot & [b] \\ \cdot & \cdot & \cdot & \cdot \end{pmatrix} \otimes \begin{pmatrix} \cdot & [a] & \cdot & \cdot \\ \cdot & \cdot & [a] & \cdot \\ [a] & \cdot & \cdot & [b] \\ \cdot & \cdot & [b] & \cdot \end{pmatrix} =$$

$$= \begin{pmatrix} \cdot & \cdot & \cdot & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \end{pmatrix}$$

### 8.3 Алгоритм

Идея алгоритма основана на обобщении пересечения двух конечных автоматов до пересечения рекурсивного автомата, построенного по грамматике, со входным графом.

Пересечение двух конечных автоматов — тензорное произведение соответствующих графов. Пересечение языков коммутативно, тензорное произведение нет, но, как было сказано в замечании 1.7.1, существует решение этой проблемы.

Будем рассматривать два конечных автомата: один получен из входного графа, второй из грамматики. Можно найти их пересечение, вычислив тензорное произведение матриц смежности соответствующих графов. Однако, одной такой итерации не достаточно для решения исходной задачи. За первую итерацию мы найдём только те пути, которые выводятся в грамматике за один шаг. После этого необходимо добавить соответствующие рёбра во входной граф и повторить операцию: так мы найдём пути, выводимые за два шага. Данные действия надо повторять до тех пор, пока не перестанут находиться новые пары достижимых вершин. Псевдокод, описывающий данные действия, представлен в листинге 7.

Алгоритм выполняется до тех пор, пока матрица смежности  $M_2$  изменяется. На каждой итерации цикла алгоритм последовательно выполняет следующие команды: пересечение двух автоматов через тензорное произведение, транзитивное замыкание результата тензорного произведения и итерация по всем ячейкам получившейся после транзитивного замыкания матрицы, что необходимо для поиска новых пар достижимых вершин. Во время итерации по ячейкам матрицы транзитивного замыкания алгоритм сначала проверяет наличие ребра в данной ячейке, а затем — принадлежность стартовой и конечной вершин ребра к стартовому и конечному состоянию входного рекурсивного автомата. При удовлетворении этих условий алгоритм добавляет нетерминал (или нетерминалы), соответствующие стартовой и конечной вершинам ребра, в ячейку матрицы  $M_2$ , полученной с помощью функции *getCoordinates(i, j)*.

Представленный алгоритм не требует преобразования грамматики в ОНФХ, более того, рекурсивный автомат может быть минимизирован. Однако, результатом тензорного произведения является матрица существенно большего размера, чем в алгоритме, основанном на матричном произведении. Кроме этого, необходимо искать транзитивное замыкание этой матрицы.

Ещё одним важным свойством представленного алгоритма является его оптимальность при обработке регулярных запросов. Так как по контекстно-свободной грамматике мы не можем поределить, задаёт ли она регулярный язык, то при добавлении в язык запросов возможности задавать контекстно-свободные ограничения, возникает проблема: мы не можем в общем случае отличить регулярный запрос от контекстно-свободного. Следовательно, мы вынуждены применять наиболее общий механизм выполнения запросов, что может приводить к существенным накладным расходам при выполнении ре-

**Listing 7** Поиск путей через тензорное произведение

---

```

1: function CONTEXTFREEPATHQUERYINGTP( $G, \mathcal{G}$ )
2:    $R \leftarrow$  рекурсивный автомат для  $G$ 
3:    $N \leftarrow$  нетерминальный алфавит для  $R$ 
4:    $S \leftarrow$  стартовые состояния для  $R$ 
5:    $F \leftarrow$  конечные состояния для  $R$ 
6:    $M_1 \leftarrow$  матрица смежности  $R$ 
7:    $M_2 \leftarrow$  матрица смежности  $\mathcal{G}$ 
8:   for  $N_i \in N$  do
9:     if  $N_i \xrightarrow{*} \varepsilon$  then
10:      for all  $j \in \mathcal{G}.V : M_2[j, j] \leftarrow M_2[j, j] \cup \{N_i\}$   $\triangleright$  Добавим петли для
      нетерминалов, выводящих  $\varepsilon$ 
11:   while матрица  $M_2$  изменяется do
12:      $M_3 \leftarrow M_1 \otimes M_2$   $\triangleright$  Пересечение графов
13:      $tC_3 \leftarrow \text{transitiveClosure}(M_3)$ 
14:      $n \leftarrow$  количество строк и столбцов матрицы  $M_3$   $\triangleright$  размер матрицы
      $M_3 = n \times n$ 
15:     for  $i \in 0..n$  do
16:       for  $j \in 0..n$  do
17:         if  $tC_3[i, j]$  then
18:            $s \leftarrow$  стартовая вершина ребра  $tC_3[i, j]$ 
19:            $f \leftarrow$  конечная вершина ребра  $tC_3[i, j]$ 
20:           if  $s \in S$  and  $f \in F$  then
21:              $x, y \leftarrow \text{getCoordinates}(i, j)$ 
22:              $M_2[x, y] \leftarrow M_2[x, y] \cup \{\text{getNonterminals}(s, f)\}$ 
23:   return  $M_2$ 

```

---



гулярного запроса. Данный же алгоритм не выполнит лишних действий, так как сразу выполнит классическое пересечение двух автоматов и получит результат.

## 8.4 Примеры

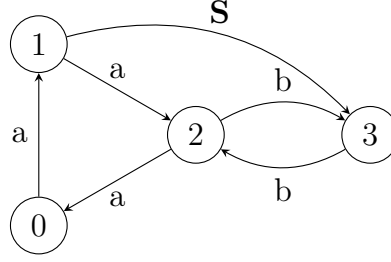
Рассмотрим подробно ряд примеров работы описанного алгоритма. Будем для каждой итерации внешнего цикла выписывать результаты основных операций: тензорного произведения, транзитивного замыкания, обновления матрицы смежности входного графа. Новые, по сравнению с предыдущим состоянием, элементы матриц будем выделять.

**Пример 8.4.1.** Теоретически худший случай. Такой же как и для матричного.

**Итерация 1 (конец).** Начало в разделе 8.2, где мы вычислили тензорное произведение матриц смежности. Теперь нам осталось только вычислить транзитивное замыкание полученной матрицы:

$$tc(M_3) = \left( \begin{array}{cccc|cccc|cccc|cccc} \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \hline \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \hline \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \hline \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \end{array} \right).$$

Мы видим, что в результате транзитивного замыкания появилось новое ребро с меткой  $ab$  из вершины  $(0, 1)$  в вершину  $(3, 3)$ . Так как вершина 0 является стартовой в рекурсивном автомате, а 3 является финальной, то слово вдоль пути из вершины 1 в вершину 3 во входном графе выводимо из нетерминала  $S$ . Это означает, что в графе должно быть добавлено ребро из 0 в 3 с меткой  $S$ , после чего граф будет выглядеть следующим образом:



Матрица смежности обновлённого графа:

$$M_2 = \begin{pmatrix} \cdot & [a] & \cdot & \cdot \\ \cdot & \cdot & [a] & [\mathbf{S}] \\ [a] & \cdot & \cdot & [b] \\ \cdot & \cdot & [b] & \cdot \end{pmatrix}$$

Итерация закончена. Возвращаемся к началу цикла и вновь вычисляем тензорное произведение.

**Итерация 2.** Вычисляем тензорное произведение матриц смежности.

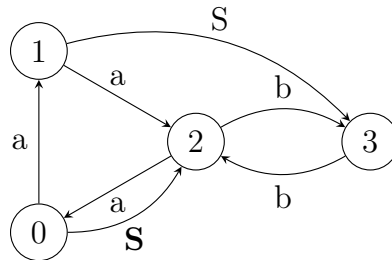
$$M_3 = M_1 \otimes M_2 = \begin{pmatrix} \cdot & [a] & \cdot & \cdot \\ \cdot & \cdot & [\mathbf{S}] & [b] \\ \cdot & \cdot & \cdot & [b] \\ \cdot & \cdot & \cdot & \cdot \end{pmatrix} \otimes \begin{pmatrix} \cdot & [a] & \cdot & \cdot \\ \cdot & \cdot & [a] & [\mathbf{S}] \\ [a] & \cdot & \cdot & [b] \\ \cdot & \cdot & [b] & \cdot \end{pmatrix} =$$

$$= \begin{pmatrix} \cdot & \cdot & \cdot & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \end{pmatrix}$$

Вычисляем транзитивное замыкание полученной матрицы:

$$tc(M_3) = \left( \begin{array}{c|c|c|c} \dots & \cdot & [a] & \cdot & \dots & [aS] & \cdot & [aSb] & \cdot \\ \dots & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & [ab] \\ \dots & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \hline \dots & \cdot & \cdot & \cdot & \cdot & \cdot & [S] & \cdot & [Sb] & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [b] & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [b] & \cdot \\ \hline \dots & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [b] \\ \dots & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [b] & \cdot \\ \hline \dots & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \end{array} \right)$$

В транзитивном замыкании появилось три новых ребра, однако только одно из них соединяет вершины, соответствующие стартовому и конечному состоянию входного рекурсивного автомата. Таким образом только это ребро должно быть добавлено во входной граф. В итоге, обновлённый граф:



И его матрица смежности:

$$M_2 = \begin{pmatrix} \cdot & [a] & [S] & \cdot \\ \cdot & \cdot & [a] & [S] \\ [a] & \cdot & \cdot & [b] \\ \cdot & \cdot & [b] & \cdot \end{pmatrix}$$

**Итерация 3.** Снова начинаем с тензорного произведения.



Матрица смежности обновлённого графа:

$$M_2 = \begin{pmatrix} \cdot & [a] & [S] & \cdot \\ \cdot & \cdot & [a] & [S] \\ [a] & \cdot & \cdot & [b, \mathbf{S}] \\ \cdot & \cdot & [b] & \cdot \end{pmatrix}$$

Уже можно заметить закономерность: на каждой итерации мы добавляем ровно одно новое ребро во входной граф, ровно как и в алгоритме, основанном на матричном произведении, и как в алгоритме Хеллингса. То есть находим ровно одну новую пару вещей, между которыми существует интересующий нас путь. Попробуйте спрогнозировать, сколько итераций нам ещё осталось сделать.

**Итерауия 4.** Продолжаем. Вычисляем тензорное произведение.

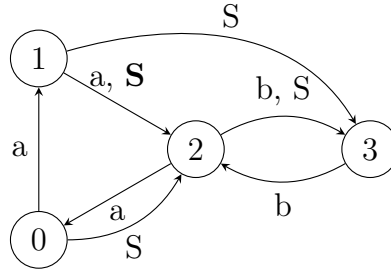
$$M_3 = M_1 \otimes M_2 = \begin{pmatrix} \cdot & [a] & \cdot & \cdot \\ \cdot & \cdot & [S] & [b] \\ \cdot & \cdot & \cdot & [b] \\ \cdot & \cdot & \cdot & \cdot \end{pmatrix} \otimes \begin{pmatrix} \cdot & [a] & [S] & \cdot \\ [a] & \cdot & [a] & [S] \\ \cdot & \cdot & [b, S] & \cdot \\ \cdot & \cdot & [b] & \cdot \end{pmatrix} =$$

$$= \begin{pmatrix} \cdot & \cdot & \cdot & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \end{pmatrix}$$

Затем транзитивное замыкание:

$$tc(M_3) = \left( \begin{array}{c|c|c|c|c} \dots & \cdot & [a] & \cdot & \cdot \\ \dots & \cdot & \cdot & [a] & \cdot \\ \dots & [a] & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot \\ \hline \dots & \cdot & \cdot & \cdot & [S] \\ \dots & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & [S] \\ \dots & \cdot & \cdot & \cdot & \cdot \\ \hline \dots & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot \\ \hline \dots & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot \\ \hline \dots & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot \\ \dots & \cdot & \cdot & \cdot & \cdot \end{array} \right)$$

И снова обновляем граф, так как в транзитивном замыкании появился один (и снова ровно один) новый элемент, соответствующий принимающему пути в автомате.

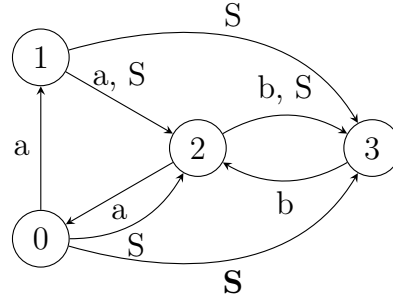


Матрица смежности обновлённого графа:

$$M_2 = \begin{pmatrix} \cdot & [a] & [S] & \cdot \\ \cdot & \cdot & [a, \mathbf{S}] & [S] \\ [a] & \cdot & \cdot & [b, S] \\ \cdot & \cdot & [b] & \cdot \end{pmatrix}$$

**Итерация 5.** Приступаем к выполнению следующей итерации основного цикла. Вычисляем тензорное произведение.





Матрица смежности обновлённого графа:

$$M_2 = \begin{pmatrix} \cdot & [a] & [S] & [\mathbf{S}] \\ \cdot & \cdot & [a, S] & [S] \\ [a] & \cdot & \cdot & [b, S] \\ \cdot & \cdot & [b] & \cdot \end{pmatrix}$$

**Итерация 6.** И наконец последняя содержательная итерация основного цикла.

$$M_3 = M_1 \otimes M_2 = \begin{pmatrix} \cdot & [a] & \cdot & \cdot \\ \cdot & \cdot & [S] & [b] \\ \cdot & \cdot & \cdot & [b] \\ \cdot & \cdot & \cdot & \cdot \end{pmatrix} \otimes \begin{pmatrix} \cdot & [a] & [S] & [S] \\ \cdot & \cdot & [a, S] & [S] \\ [a] & \cdot & \cdot & [b, S] \\ \cdot & \cdot & [b] & \cdot \end{pmatrix} =$$

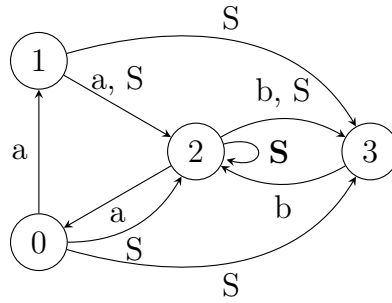
$$= \begin{pmatrix} \cdot & \cdot & \cdot & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \end{pmatrix}$$

Транзитивное замыкание:



$$tc(M_3) = \left( \begin{array}{cccc|cccc|cccc|cccc} \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \end{array} \right)$$

Обновлённый граф:



И матрица смежности:

$$M_2 = \begin{pmatrix} \cdot & [a] & [S] & [S] \\ \cdot & \cdot & [a, S] & [S] \\ [a] & \cdot & \mathbf{[S]} & [b, S] \\ \cdot & \cdot & [b] & \cdot \end{pmatrix}$$

Следующая итерация не приведёт к изменению графа. Читатель может убедиться в этом самостоятельно. Соответственно, алгоритм можно завершать. Нам потребовалось семь итераций (шесть содержательных и одна проверочная), на каждой из которых вычисляются тензорное произведение двух матриц смежности и транзитивное замыкание результата.

Матрица смежности получилась такая же, как и раньше, ответ правильный. Мы видим, что количество итераций внешнего цикла такое же как и у алгоритма СУК (пример 6.2.1).

**Пример 8.4.2.** В данном примере мы увидим, как структура грамматики и, соответственно, рекурсивного автомата, влияет на процесс вычислений.

Интуитивно понятно, что чем меньше состояний в рекурсивной сети, тем лучше. То есть желательно получить как можно более компактное описание контекстно-свободного языка.

Для примера возьмём в качестве КС языка язык Дика на одном типе скобок и опишем его двумя различными грамматиками. Первая грамматика классическая:

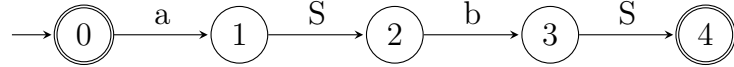
$$G_1 = \langle \{a, b\}, \{S\}, \{S \rightarrow a S b S \mid \varepsilon\} \rangle$$

Во второй грамматике мы будем использовать конструкции регулярных выражений в правой части правил. То есть вторая грамматика находится в EBNF (Расширенная форма Бэкуса-Наура [33, 73]).

$$G_2 = \langle \{a, b\}, \{S\}, \{S \rightarrow (a S b)^*\} \rangle$$

Построим рекурсивные автоматы  $N_1$  и  $N_2$  и их матрицы смежности для этих грамматик.

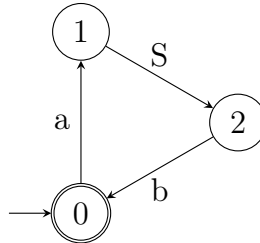
Рекурсивный автомат  $N_1$  для грамматики  $G_1$ :



Матрица смежности  $N_1$ :

$$M_1^1 = \begin{pmatrix} \cdot & [a] & \cdot & \cdot & \cdot \\ \cdot & \cdot & [S] & \cdot & \cdot \\ \cdot & \cdot & \cdot & [b] & \cdot \\ \cdot & \cdot & \cdot & \cdot & [S] \\ \cdot & \cdot & \cdot & \cdot & \cdot \end{pmatrix}$$

Рекурсивный автомат  $N_2$  для грамматики  $G_2$ :

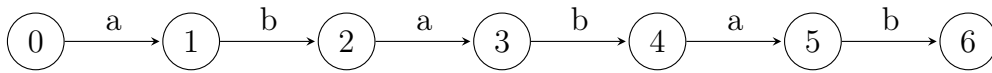


Матрица смежности  $N_2$ :

$$M_1^2 = \begin{pmatrix} \cdot & [a] & \cdot \\ \cdot & \cdot & [S] \\ [b] & \cdot & \cdot \end{pmatrix}$$

Первое очевидное наблюдение — количество состояний в  $N_2$  меньше, чем в  $N_1$ . Это значит, что матрица смежности, а значит, и результат тензорного произведения будут меньше, следовательно, вычисления будут производиться быстрее.

Выполним пошагово алгоритм для двух заданных рекурсивных сетей на линейном входе:



Сразу дополним матрицу смежности нетерминалами, выводящими пустую строку, и получим следующую матрицу:

$$M_2 = \begin{pmatrix} [S] & [a] & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & [S] & [b] & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & [S] & [a] & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & [S] & [b] & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & [S] & [a] & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & [S] & [b] \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [S] \end{pmatrix}$$

Сперва запустим алгоритм на данном входе и  $N_2$ . Первый шаг первой итерации — вычисление тензорного произведения  $M_1^2 \otimes M_2$ .

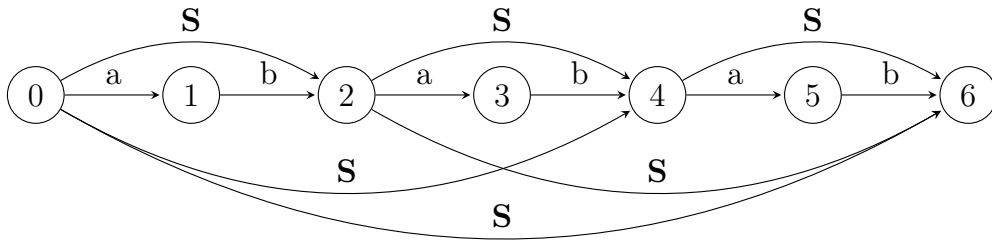
[illegible]

Опустим промежуточные шаги вычисления транзитивного замыкания  $M_3$  и сразу представим конечный результат:

$$tc(M_3) = \left( \begin{array}{ccc|ccc|cccccc} \text{[aSb]} & \text{[aSbaSb]} & \text{[aSbaSbaSb]} & [a] & \text{[aSba]} & \text{[aSbaSba]} & \text{[aS]} & \text{[aSbaS]} & \text{[aSbaSbaS]} & \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \text{[aSb]} & \text{[aSbaSb]} & \cdot & [a] & \text{[aSba]} & \cdot & \cdot & \text{[aS]} & \text{[aSbaS]} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \text{[aSb]} & \cdot & \cdot & [a] & \cdot & \cdot & \cdot & \text{[aS]} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \hline \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [S] & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [S] & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [S] & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [S] & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [S] & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [S] \\ \hline [b] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & [b] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & [b] & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \end{array} \right)$$

В результате вычисления транзитивного замыкания появилось большое количество новых рёбер, однако нас будут интересовать только те, информация о которых храниться в левом верхнем блоке. Остальные рёбра не соответствуют принимающим путям в рекурсивном автомате (убедитесь в этом самостоятельно).

После добавления соответствующих рёбер, мы получим следующий граф:



Его матрица смежности:

$$M_2 = \begin{pmatrix} [S] & [a] & [S] & \cdot & [S] & \cdot & [S] \\ \cdot & [S] & [b] & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & [S] & [a] & [S] & \cdot & [S] \\ \cdot & \cdot & \cdot & [S] & [b] & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & [S] & [a] & [S] \\ \cdot & \cdot & \cdot & \cdot & \cdot & [S] & [b] \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & [S] \end{pmatrix}$$

Таким образом видно, что для выбранных входных данных алгоритму достаточно двух итераций основного цикла: первая содержательная, вторая, как обычно, проверочная. Читателю предлагается самостоятельно выяснить, сколько умножений матриц потребуется, чтобы вычислить транзитивное замыкание на первой итерации.

Теперь запустим алгоритм на второй грамматике и том же входе.



$$tc(M_3) = \left( \begin{array}{cccc} \begin{array}{c} [a] \\ [a] \\ [a] \end{array} & \begin{array}{c} [aS] \\ [aS] \\ [aS] \end{array} & \begin{array}{c} [aSb] \\ [aSb] \\ [aSb] \end{array} & \begin{array}{c} [aSbS] \\ [aSbS] \\ [aSbS] \end{array} \\ \begin{array}{c} [S] \\ [S] \\ [S] \\ [S] \\ [S] \end{array} & \begin{array}{c} [S] \\ [S] \\ [S] \\ [S] \\ [S] \end{array} & \begin{array}{c} [b] \\ [b] \\ [b] \end{array} & \begin{array}{c} [S] \\ [S] \\ [S] \\ [S] \\ [S] \end{array} \\ \begin{array}{c} [b] \\ [b] \\ [b] \end{array} & \begin{array}{c} [b] \\ [b] \\ [b] \end{array} & \begin{array}{c} [b] \\ [b] \\ [b] \end{array} & \begin{array}{c} [S] \\ [S] \\ [S] \\ [S] \\ [S] \end{array} \end{array} \right)$$

$$M_2 = \begin{pmatrix} [S] & [a] & \mathbf{[S]} & \cdot & \cdot & \cdot & \cdot \\ \cdot & [S] & [b] & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & [S] & [a] & \mathbf{[S]} & \cdot & \cdot \\ \cdot & \cdot & [S] & [b] & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & [S] & [a] & \mathbf{[S]} & \cdot \\ \cdot & \cdot & \cdot & \cdot & [S] & [b] & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & [S] & \cdot \end{pmatrix}$$

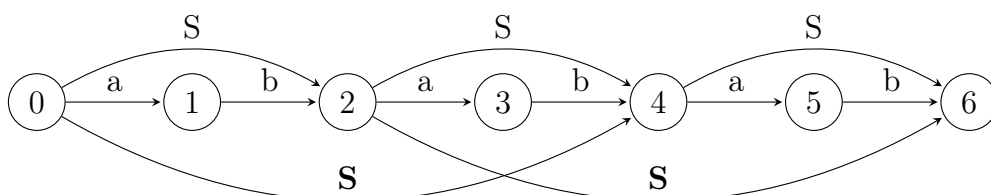
Потребуется ещё одна итерация.





$$tc(M_3) = \left( \begin{array}{cccccccc} \dots & [a] & \dots & \dots & [aS] & \dots & [aS] & \dots & [aSb] & \dots & [aSb] & \dots & [aSbS] & \dots & [aSbS] & \dots & [aSbS] & \dots \\ \dots & \dots & [a] & \dots & \dots & [aS] & \dots & [aS] & \dots & [aSb] & \dots & [aSb] & \dots & [aSbS] & \dots & [aSbS] & \dots & [aSbS] & \dots \\ \dots & \dots & \dots & [a] & \dots & \dots & [aS] & \dots & \dots & [aSb] & \dots & \dots & \dots & [aSbS] & \dots & \dots & \dots & [aSbS] & \dots \\ \dots & \dots & \dots & \dots & [S] & [S] & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & [S] & [S] & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & [S] & [S] & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & [S] & [S] & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & [S] & [S] & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & [S] & [S] & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & [S] & [S] & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & [S] & [S] & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & [S] & [S] & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & [S] & [S] & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & [S] & [S] & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & [S] & [S] & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & [S] & [S] & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & [S] & [S] \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & [S] \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \end{array} \right)$$

Обновлённый граф:



На этом шаге мы смогли “склеить” из подстрок, выводимых из  $S$ , более длинные пути. Однако, согласно правилам грамматики, мы смогли “склеить” только две подстроки в единое целое.

Матрица смежности обновлённого графа:

$$M_2 = \begin{pmatrix} [S] & [a] & [S] & \dots & [S] & \dots \\ \dots & [S] & [b] & \dots & \dots & \dots \\ \dots & \dots & [S] & [a] & [S] & \dots \\ \dots & \dots & \dots & [S] & [b] & \dots \\ \dots & \dots & \dots & \dots & [S] & [a] & [S] \\ \dots & \dots & \dots & \dots & \dots & [S] & [b] \\ \dots & \dots & \dots & \dots & \dots & \dots & [S] \end{pmatrix}$$

И, наконец, последняя содержательная итерация.

[illegible]

Транзитивное замыкание:

[illegible]

В конечном итоге мы получаем такой же результат, как и при первом запуске. Однако нам потребовалось выполнить существенно больше итераций внешнего цикла, а именно четыре (три содержательных и одна проверочная). При этом, в ходе работы нам пришлось оперировать существенно большими матрицами:  $35 \times 35$  против  $21 \times 21$ .

Таким образом, видно, что минимизация представления запроса, в частности, минимизация рекурсивного автомата как конечного автомата над смешанным алфавитом может улучшить производительность выполнения запросов.

## 8.5 Особенности реализации

Как и алгоритмы, представленные в разделе 7, представленный здесь алгоритм оперирует разреженными матрицами, поэтому, к нему применимы все те же соображения, что и к алгоритмам, основанным на произведении матриц. Более того, так как результат тензорного произведения является блочной матрицей, то могут оказаться полезными различные форматы для хранения

ния блочно-разреженных матриц. Вместе с этим, в некоторых случаях матрицу смежности рекурсивного автомата удобнее представлять в классическом, плотном, виде, так как для некоторых запросов её размер мал и накладные расходы на представление в разреженном формате и работе с ним будут больше, чем выигрыш от его использования.

Также заметим, что блочная структура матриц даёт хорошую основу для распределённого умножения матриц при построении транзитивного замыкания.

Вместо того, чтобы перезаписывать каждый раз матрицу смежности входного графа  $M_2$  можно вычислять только разницу с предыдущим шагом. Для этого, правда, потребуется хранить в памяти ещё одну матрицу. Поэтому нужно проверять, что вычислительно дешевле: поддерживать разницу и потом каждый раз поэлементно складывать две матрицы или каждый раз вычислять полностью произведение.

Заметим, что для решения задачи достижимости нам не нужно накапливать пути вдоль рёбер, как мы это делали в примерах, соответственно, во-первых, можно переопределить тензорное произведение так, чтобы его результатом являлась булева матрица, во-вторых, как следствие первого изменения, транзитивное замыкание для булевой матрицы можно искать с применением соответствующих оптимизаций.

## Глава 9

# Сжатое представление леса разбора

Матричный алгоритм даёт нам ответ на вопрос о достижимости, но не предоставляет самих путей. Что делать, если мы хотим построить все пути, удовлетворяющие ограничениям?

Проблема в том, что искомое множество путей может быть бесконечным. Можем ли мы предложить конечную структуру, однозначно описывающую такое множество? Вспомним, что пересечение контекстно-свободного языка с регулярным — это контекстно-свободный язык. Мы знаем, что контекстно-свободный язык можно описать контекстно-свободной грамматикой, которая конечна. Это и есть решение нашего вопроса. Осталось только научиться строить такую грамматику.

Прежде, чем двинуться дальше, рекомендуется вспомнить всё, что касается деревьев вывода 5.1.

### 9.1 Лес разбора как представление контекстно-свободной грамматики

Для начала нам потребуется внести некоторые изменения в конструкцию дерева вывода.

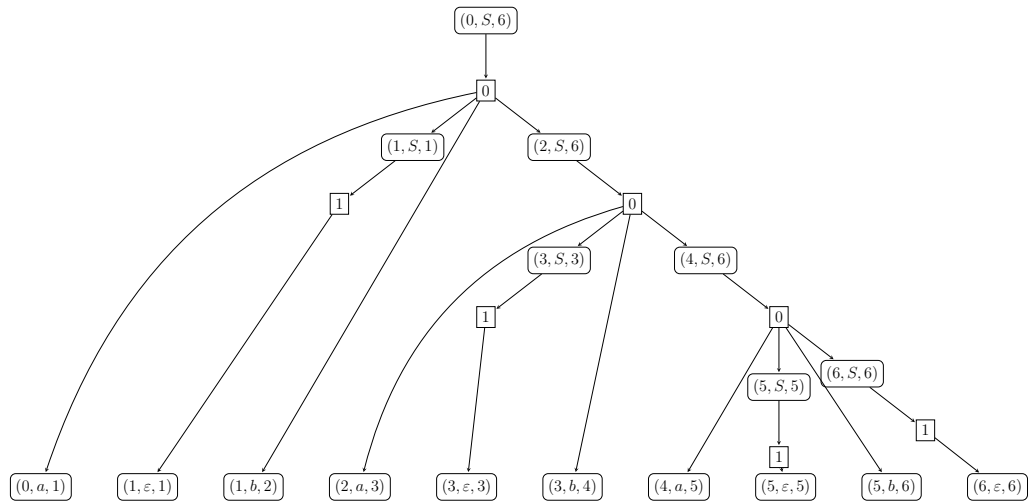
Во-первых, заметим, что в дереве вывода каждый узел соответствует выводу какой-то подстроки с известными позициями начала и конца. Давайте будем сохранять эту информацию в узлах дерева. Таким образом, метка любого узла это тройка вида  $(i, q, j)$ , где  $i$  — координата начала подстроки, со-

ответствующей этому узлу,  $j$  — координата конца,  $q \in \Sigma \cup N$  — метка как в исходном определении.

Во-вторых, заметим, что внутренний узел со своими сыновьями связаны с продукцией в грамматике: узел появляется благодаря применению конкретной продукции в процессе вывода. Давайте занумеруем все продукции в грамматике и добавим в дерево вывода ещё один тип узлов (дополнительные узлы), в которых будем хранить номер применённой продукции. Получим следующую конструкцию: непосредственный предок дополнительного узла — это левая часть продукции, а непосредственные сыновья дополнительного узла — это правая часть продукции.

**Пример 9.1.1.** Построим модифицированное дерево вывода цепочки  $0a_1b_2a_3b_4a_5b_6$  в грамматике

$$G_0 = \langle \{a, b\}, \{S\}, S, \{ \begin{array}{l} (0)S \rightarrow a S b S, \\ (1)S \rightarrow \varepsilon \end{array} \} \rangle$$



Сохраняемая нами дополнительная информация позволит переиспользовать узлы в том случае, если деревьев вывода оказалось несколько (в случае неоднозначной грамматики). При этом мы можем не бояться, что переиспользование узлов может привести к появлению ранее несуществовавших деревьев вывода, так как дополнительная информация позволяет делать только “безопасные” склейки и затем восстанавливать только корректные деревья. Таким

образом, мы можем представить лес вывода в виде единой структуры данных без дублирования информации.

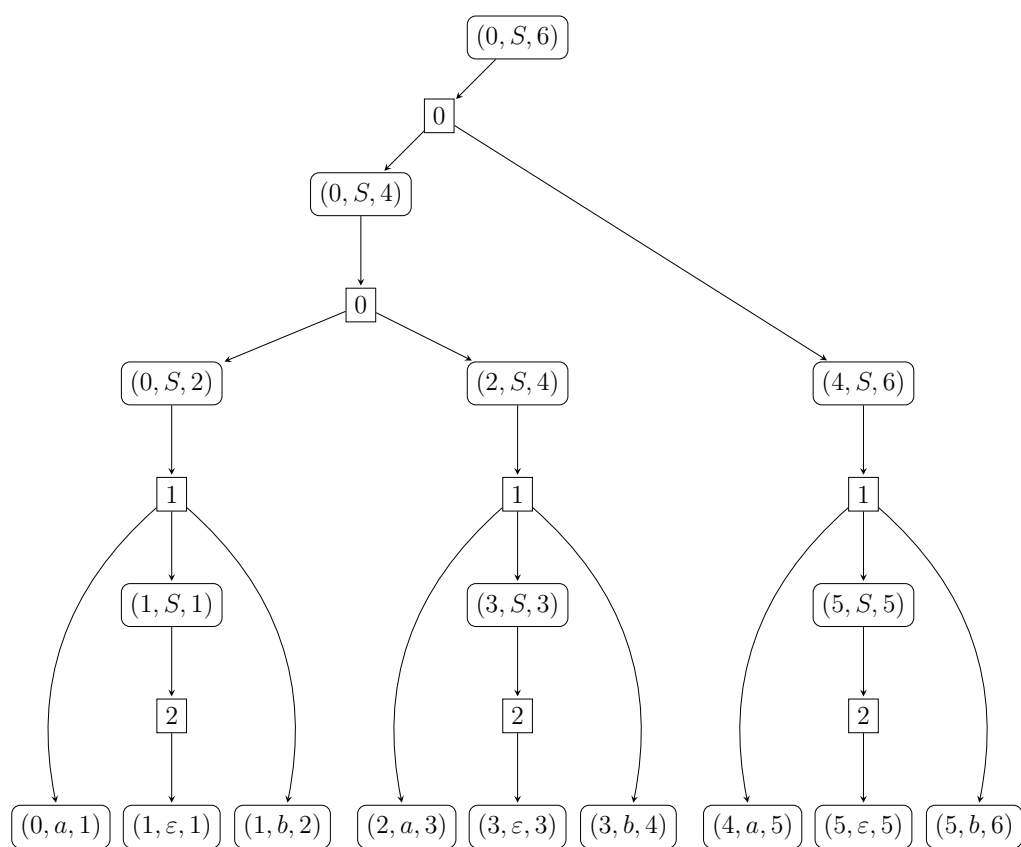
**Пример 9.1.2.** Сжатие леса вывода. Построим несколько деревьев вывода цепочки  $a_1b_2a_3b_4a_5b_6$  в грамматике

$$G_1 = \langle \{a, b\}, \{S\}, S, \{ \begin{array}{l} (0) S \rightarrow SS, \\ (1) S \rightarrow a S b, \\ (2) S \rightarrow \varepsilon \end{array} \rangle$$

Предположим, что мы строим левосторонний вывод. Тогда после первого применения продукции 0 у нас есть два варианта переписывания первого нетерминала: либо с применением продукции 0, либо с применением продукции 1:

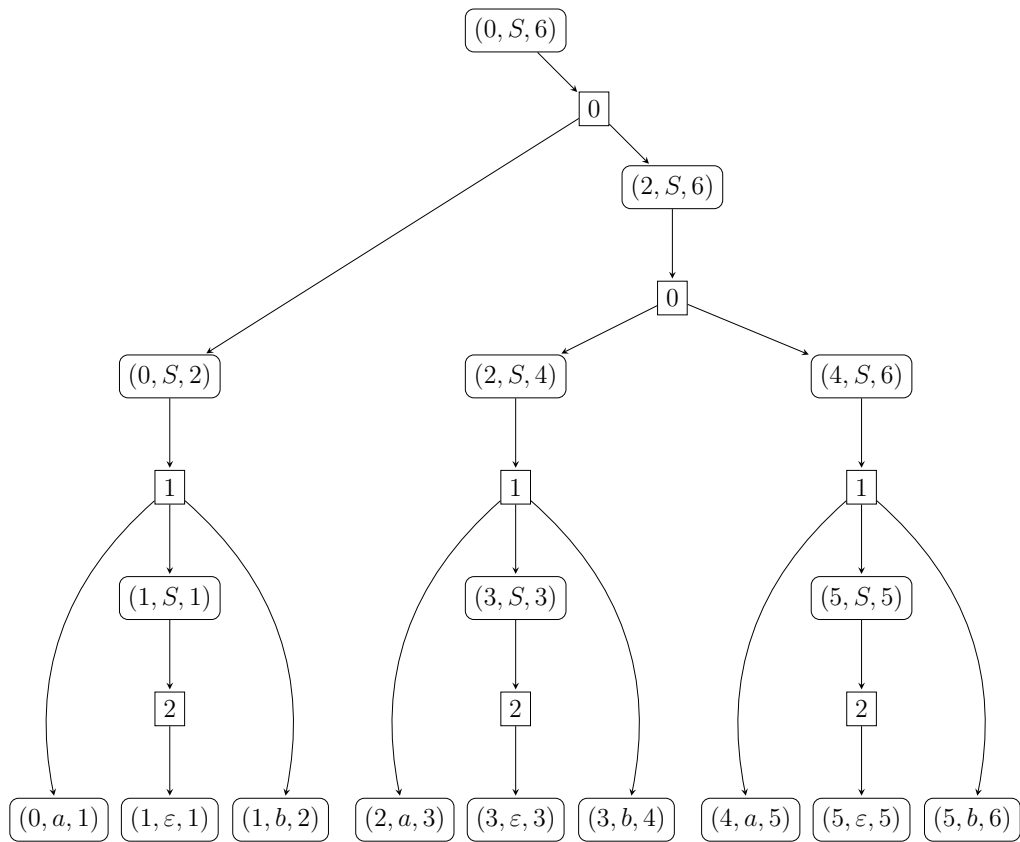
$$\begin{array}{l} S \xrightarrow{0} SS \xrightarrow{0} SSS \xrightarrow{1} aSbSS \xrightarrow{2} abSS \xrightarrow{1} abaSbS \xrightarrow{2} ababS \xrightarrow{1} ababaSb \xrightarrow{2} ababab \\ S \xrightarrow{0} SS \xrightarrow{1} aSbS \xrightarrow{2} abS \xrightarrow{0} abSS \xrightarrow{1} abaSbS \xrightarrow{2} ababS \xrightarrow{1} ababaSb \xrightarrow{2} ababab \end{array}$$

Сначала рассмотрим первый вариант (применили переписывание по продукции 0). Все остальные шаги вывода деретерминированы и в результате мы получим следующее дерево разбора:

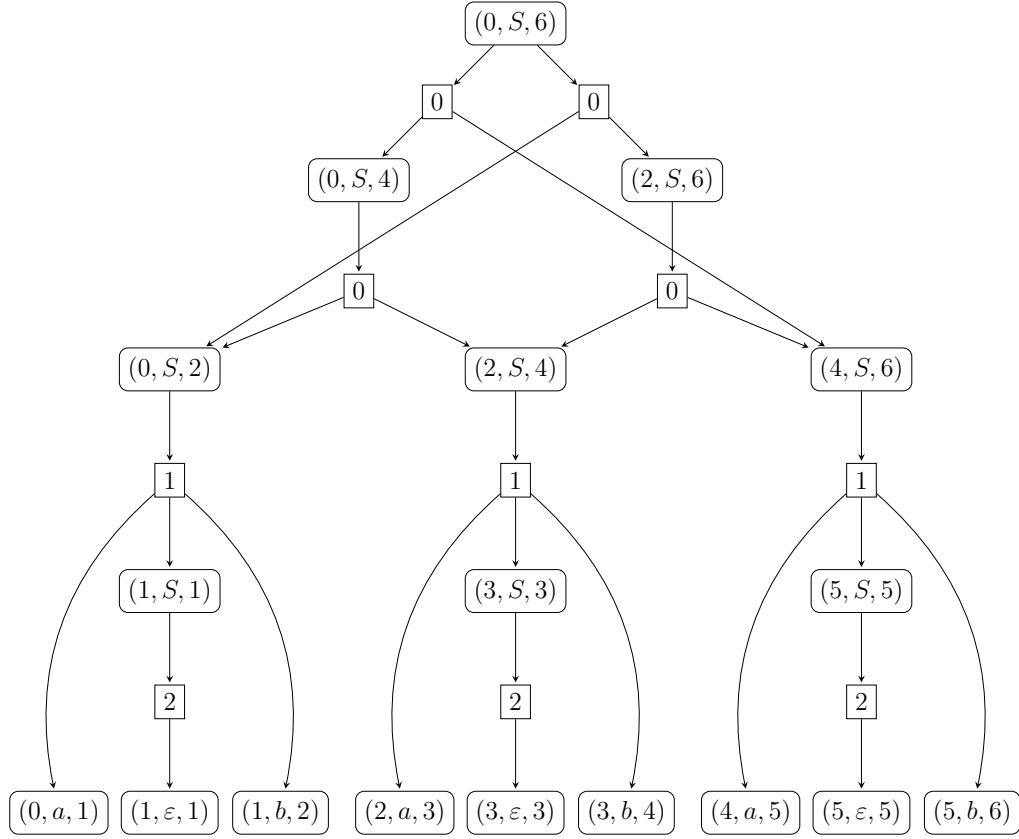


Теперь рассмотрим второй вариант — применить продукцию 1. Остальные шаги вывода всё также детерминированы. В результате мы получим следующее дерево вывода:





В двух построенных деревьях большое количество одинаковых узлов. Построим структуру, которая содержит оба дерева и при этом никакие нетерминальные и терминальные узлы не встречаются дважды. В результате мы получим следующий граф:



Мы получили очень простой вариант сжатого представления леса разбора (Shared Packed Parse Forest, SPPF). Впервые подобная идея была предложена Джоаном Рекерсом в его кандидатской диссертации [50]. В дальнейшем она нашла широкое применение в обобщённом (generalized) синтаксическом анализе и получила серьёзное развитие. В частности, наш вариант, хоть и позволяет избежать экспоненциального разрастания леса разбора, всё же не является оптимальным. Оптимальное асимптотическое поведение достигается при использовании бинаризованного SPPF [14] — в этом случае объём леса составляет  $O(n^3)$ , где  $n$  — это длина входной строки.

Различные модификации SPPF применяются в таких алгоритмах синтаксического анализа, как RNGLR [55], бинаризованная версия SPPF в BRNGLR [58] и GLL [56, 1]<sup>1</sup>.

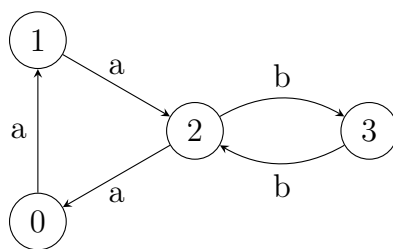
В действительности SPPF может содержать в себе циклы. Для линейного входа их можно получить, когда есть возможность выводить по грамматике бесконечные эпсилон-цепочки. Циклы будут вырожденными, но они будут.

<sup>1</sup>Ещё немного полезной информации про SPPF: <http://www.bramvandersanden.com/post/2014/06/shared-packed-parse-forest/>.

Мы, кроме традиционного использования, будем применять SPPF для представления результатов КС запросов к графам.

В графе может существовать множество способов получить путь из одной вершины в другую. И точно так же при построении деревьев вывода путей может появиться несколько одинаковых нетерминалов, получаемых в разных деревьях по-разному. При объединении в SPPF может оказаться, что какой-то путь из вершины  $a$  в вершину  $b$  является подпутем другого пути из вершины  $a$  в вершину  $b$ , просто более длинного. То есть появятся циклические зависимости.

**Пример 9.1.3.** Рассмотрим пример SPPF для задачи поиска путей с КС ограничениями. Пусть дан граф  $\mathcal{G}$  :



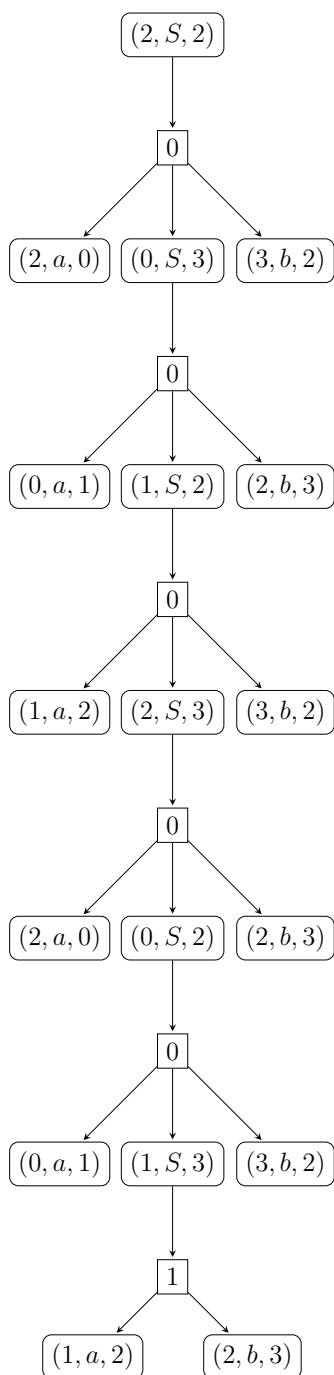
Дана грамматика

$$G = \langle \{a, b\}, \{S\}, S, \{ \begin{array}{l} (0) S \rightarrow a S b, \\ (1) S \rightarrow a b, \end{array} \} \rangle$$

Попробуем найти все пути из вершины 2 в вершину 2, выводимые из нетерминала  $S$ . Проверить наличие такого пути можно используя уже известные нам алгоритмы, однако сами пути пока будем строить “методом пристального взгляда”. Найдем один из них. Пусть это будет

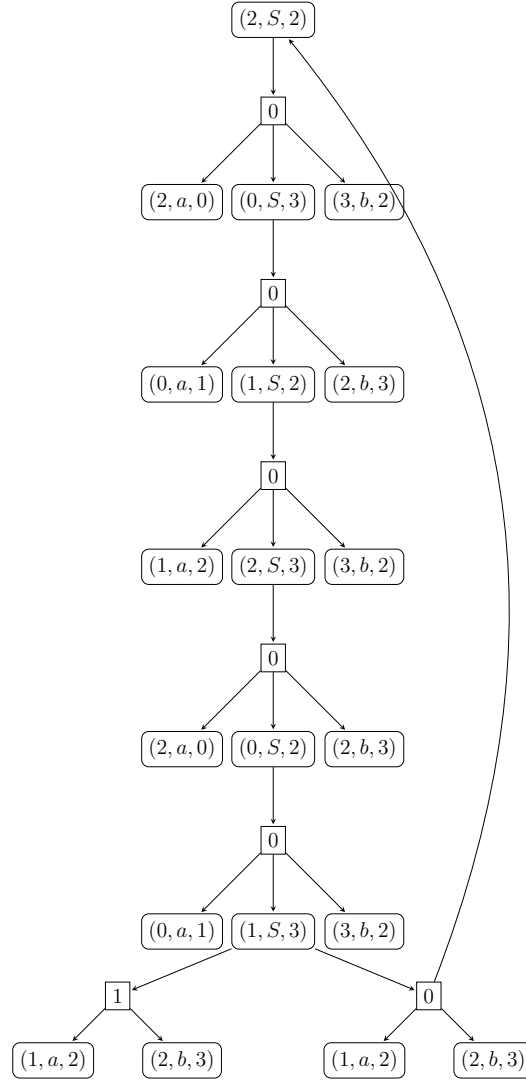
$$2 \xrightarrow{a} 0 \xrightarrow{a} 1 \xrightarrow{a} 2 \xrightarrow{a} 0 \xrightarrow{a} 1 \xrightarrow{a} 2 \xrightarrow{b} 3 \xrightarrow{b} 2 \xrightarrow{b} 3 \xrightarrow{b} 2 \xrightarrow{b} 3 \xrightarrow{b} 2.$$

Построим дерево его вывода.



Мы построили дерево вывода для одного пути из вершины 2 в неё же. Но можно заметить, что таких путей бесконечно много: мы можем бесконечное число раз повторять уже выполненный обход и получать всё более длинные пути. В терминах дерева вывода это будет означать, что к узлу  ${}_1S_3$  мы добавим сына, соответствующего применению продукции 0, а не 1 для нетерминала  $S$ .

В таком случае мы получим узел  ${}_2S_2$ , который уже существует в дереве и таким образом замкнём цикл.



Таким образом мы построили SPPF. Обойдя эту структуру необходимое количество раз, мы можем получить любой путь, удовлетворяющий условию. Более того, в полученном графе можно получать любые другие пути по соответствующим нетерминалам и парам вершин, содержащимся в узлах леса.

**Утверждение 9.1.1.** SPPF построенный для данной контекстно-свободной грамматики  $G$  и графа  $\mathcal{G}$

1. содержит терминальный узел вида  $(i, t_k, j)$  тогда и только тогда, когда в графе  $\mathcal{G}$  есть ребро  $(i, t_k, j)$ ;

2. содержит нетерминальный узел вида  $(i, S_k, j)$  тогда и только тогда, когда в графе  $\mathcal{G}$  есть путь из вершины  $i$  в вершину  $j$ , выводимый из нетерминала  $S_k$  в грамматике  $G$ .

Осталось увидеть, что SPPF является представлением контекстно-свободной грамматики, описывающей результат пересечения исходных графа и грамматики. Для этого просто построим грамматику  $G_{SPPF} = \langle \Sigma_{SPPF}, N_{SPPF}, S_{SPPF}, P_{SPPF} \rangle$  по SPPF следующим образом:

- $\Sigma_{SPPF}$  — все листья SPPF;
- $N_{SPPF}$  — все нетерминальные узлы SPPF;
- $S_{SPPF}$  — нетерминал, соответствующий пути, который нас будет интересовать;
- $P_{SPPF}$  — для каждого дополнительного узла (с номером продукции) добавляем продукцию, левая часть которой — непосредственный предок этого узла, а правая часть — непосредственные потомки.

**Пример 9.1.4.** Построим грамматику для полученного SPPF:

$$\begin{array}{ll}
 (0) \ {}_2S_2 \rightarrow \ {}_2a_0 \ {}_0S_3 \ {}_3b_2 & (4) \ {}_0S_2 \rightarrow \ {}_0a_1 \ {}_1S_3 \ {}_3b_2 \\
 (1) \ {}_0S_3 \rightarrow \ {}_0a_1 \ {}_1S_2 \ {}_2b_3 & (5) \ {}_1S_3 \rightarrow \ {}_1a_2 \ {}_2S_2 \ {}_2b_3 \\
 (2) \ {}_1S_2 \rightarrow \ {}_1a_2 \ {}_2S_3 \ {}_3b_2 & (6) \ {}_1S_3 \rightarrow \ {}_1a_2 \ {}_2b_3 \\
 (3) \ {}_2S_3 \rightarrow \ {}_2a_0 \ {}_0S_2 \ {}_2b_3 &
 \end{array}$$

Видим, что для одного единственного нетерминала  ${}_1S_3$  существует 2 правила, одно из которых рекурсивное. Попробуем получить левосторонний вывод

какой-нибудь цепочки в этой грамматике:

$$\begin{aligned}
& \mathbf{2S_2} \xrightarrow{(0)} \\
& {}_2a_0 \mathbf{0S_3} {}_3b_2 \xrightarrow{(1)} \\
& {}_2a_0 {}_0a_1 \mathbf{1S_2} {}_2b_3 {}_3b_2 \xrightarrow{(2)} \\
& {}_2a_0 {}_0a_1 {}_1a_2 \mathbf{2S_3} {}_3b_2 {}_2b_3 {}_3b_2 \xrightarrow{(3)} \\
& {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 \mathbf{0S_2} {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 \xrightarrow{(4)} \\
& {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 \mathbf{1S_3} {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 \xrightarrow{(5)} \\
& {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 {}_1a_2 \mathbf{2S_2} {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 \xrightarrow{(0)} \\
& {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 \mathbf{0S_3} {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 \xrightarrow{(1)} \\
& {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 \mathbf{1S_2} {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 \xrightarrow{(2)} \\
& {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 {}_1a_2 \mathbf{2S_3} {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 \xrightarrow{(3)} \\
& {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 \mathbf{0S_2} {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 \xrightarrow{(4)} \\
& {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 \mathbf{1S_3} {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 \xrightarrow{(6)} \\
& {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 {}_1a_2 {}_2a_0 {}_0a_1 {}_1a_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2 {}_2b_3 {}_3b_2
\end{aligned}$$

Мы получили цепочку, которая действительно является путем из вершины 2 в вершину 2 в заданном графе. Таким образом выводятся и любые другие соответствующие пути.





## Глава 10

# Алгоритм на основе нисходящего анализа

В данном разделе мы рассмотрим семейство алгоритмов нисходящего синтаксического (рекурсивный спуск, LL, GLL [56, 1]) рассмотрим их обобщение для задачи поиска путей с контекстно-свободными ограничениями.

### 10.1 Рекурсивный спуск

Идея рекурсивного спуска основана на использовании программного стека вызовов в качестве стека магазинного автомата. Достигается это следующим образом.

- Для каждого нетерминала создаётся функция, принимающая ещё не обработанный остаток строки и возвращающая пару: результат вывода префикса данной строки из соответствующего нетерминала и не обработанный остаток строки. В случае распознавателя результат вывода — логическое значение (выводится/не выводится).
- Каждая такая функция реализовывает обработку цепочки согласно правым частям правил для соответствующих нетерминалов: считывание символа входа при обработке терминального символа, вызов соответствующей функции при обработке нетерминального.

У такого подхода есть два ограничения:

1. Не работает с леворекурсивными грамматиками, грамматиками, в которых вывод может принимать следующий вид:

$$S \rightarrow \dots \rightarrow \underline{N_i}\alpha \rightarrow \dots \underline{N_i}\beta \rightarrow \dots \omega$$

2. Шаги должны быть однозначными.

**Пример 10.1.1.** Построим функцию рекурсивного спуска для продукции  $S \rightarrow aSbS \mid \varepsilon$ .

---

**Listing 8** Функция рекурсивного спуска

---

```

1: function S( $\omega$ )
2:   if ( $\text{len}(\omega) = 0$ ) then                                      $\triangleright$  Пустая цепочка выводима из  $S$ 
3:     return ( $\text{true}, \omega$ )
4:   if ( $\omega = a :: tl$ ) then                                        $\triangleright$  Выводимая из  $S$  подстрока должна начинаться с  $a$ 
5:      $\text{res}, tl' = \text{S}(tl)$                                             $\triangleright$  Затем должна идти подстрока, выводимая из  $S$ 
6:     if  $\text{res} \ \&\& \ tl' = b :: tl''$  then                          $\triangleright$  Если вызов закончился успешно, то надо
        проверить, что следующий символ — это  $b$ 
7:       return  $\text{S}(tl'')$                                             $\triangleright$  И снова попробовать вывести префикс из  $S$ 
8:     else
9:       return ( $\text{false}, tl'$ )
10:  else
11:    return ( $\text{false}, \omega$ )

```

---

Если возвращаемое значение этой функции — пара вида  $(\text{true}, //)$ , то разбор завершился успехом.

Данный подход применяется как для ручного написания синтаксических анализаторов, так и при генерации анализаторов по грамматике.

## 10.2 LL(k)-алгоритм синтаксического анализа

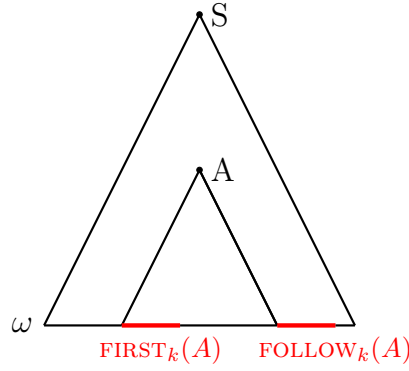
LL(k) — алгоритм синтаксического анализа — нисходящий анализ без отката, но с предпросмотром. Решение о том, какую продукцию применять, принимается на основании  $k$  следующих за текущим символом. Временная сложность алгоритма  $O(n)$ , где  $n$  — длина слова.

Алгоритм использует входной буфер, стек для хранения промежуточных данных и таблицу анализатора, которая управляет процессом разбора. В ячейке таблицы указано правило, которое нужно применять, если рассматривается

нетерминал  $A$ , а следующие  $m$  символов строки —  $t_1 \dots t_m$ , где  $m \leq k$ . Также в таблице выделена отдельная колонка для  $\$$  — маркера конца строки.

|     | ... | $t_1 \dots t_m$        | ... | $\$$ |
|-----|-----|------------------------|-----|------|
| ... | ... | ...                    | ... | ...  |
| $A$ | ... | $A \rightarrow \alpha$ | ... | ...  |
| ... | ... | ...                    | ... | ...  |

Для построения таблицы вычисляются множества  $\text{FIRST}_k$  и  $\text{FOLLOW}_k$ . Идейно их можно понимать, как первые или, соответственно, последующие  $k$  символов в результирующем выводе, при использовании нетерминала  $A$ . Данную мысль хорошо иллюстрирует рисунок:



Определим их формально:

**Определение 10.2.1.** Пусть  $G = \langle N, \Sigma, P, S \rangle$  — КС-грамматика. Множество  $\text{FIRST}_k$  определено для сентенциальной формы  $\alpha$  следующим образом:

$$\text{FIRST}_k(\alpha) = \{\omega \in \Sigma^* \mid \alpha \xRightarrow{*} \omega \text{ и } |\omega| < k \text{ либо } \exists \beta : \alpha \xRightarrow{*} \omega\beta \text{ и } |\omega| = k\}$$

, где  $\alpha, \beta \in (N \cup \Sigma)^*$ .

**Определение 10.2.2.** Пусть  $G = \langle N, \Sigma, P, S \rangle$  — КС-грамматика. Множество  $\text{FOLLOW}_k$  определено для сентенциальной формы  $\beta$  следующим образом:

$$\text{FOLLOW}_k(\beta) = \{\omega \in \Sigma^* \mid \exists \gamma, \alpha : S \xRightarrow{*} \gamma\beta\alpha \text{ и } \omega \in \text{FIRST}_k(\alpha)\}$$

В частном случае для  $k = 1$ :

$$\text{FIRST}_1(\alpha) = \{a \in \Sigma \mid \exists \gamma \in (N \cup \Sigma)^* : \alpha \xRightarrow{*} a\gamma\}, \text{ где } \alpha \in (N \cup \Sigma)^*$$

$$\text{FOLLOW}_1(\beta) = \{a \in \Sigma \mid \exists \gamma, \alpha \in (N \cup \Sigma)^* : S \xRightarrow{*} \gamma\beta a\alpha\}, \text{ где } \beta \in (N \cup \Sigma)^*$$

Множество  $\text{FIRST}_1$  можно вычислить, пользуясь следующими соотношениями:

- $\text{FIRST}_1(a\alpha) = \{a\}, a \in \Sigma, \alpha \in (N \cup \Sigma)^*$
- $\text{FIRST}_1(\varepsilon) = \{\varepsilon\}$
- $\text{FIRST}_1(\alpha\beta) = \text{FIRST}_1(\alpha) \cup (\text{FIRST}_1(\beta), \text{ если } \varepsilon \in \text{FIRST}_1(\alpha))$
- $\text{FIRST}_1(A) = \text{FIRST}_1(\alpha) \cup \text{FIRST}_1(\beta), \text{ если в грамматике есть правило } A \rightarrow \alpha \mid \beta$

Алгоритм для вычисления множества  $\text{FOLLOW}_1$ :

- Положим  $\text{FOLLOW}_1(X) = \emptyset, \forall X \in N$
- $\text{FOLLOW}_1(S) = \text{FOLLOW}_1(S) \cup \{\$, \}$ , где  $S$  — стартовый нетерминал
- Для всех правил вида  $A \rightarrow \alpha X \beta : \text{FOLLOW}_1(X) = \text{FOLLOW}_1(X) \cup (\text{FIRST}_1(\beta) \setminus \{\varepsilon\})$
- Для всех правил вида  $A \rightarrow \alpha X$  и  $A \rightarrow \alpha X \beta$ , где  $\varepsilon \in \text{FIRST}_1(\beta) : \text{FOLLOW}_1(X) = \text{FOLLOW}_1(X) \cup \text{FOLLOW}_1(A)$
- Последние два пункта применяются пока есть что добавлять в строящиеся множества.

**Пример 10.2.1.** Рассмотрим грамматику  $G$  со следующими продукциями:

$$\begin{array}{ll} S \rightarrow aS' & A' \rightarrow b \mid a \\ S' \rightarrow AbBS' \mid \varepsilon & B \rightarrow c \mid \varepsilon \\ A \rightarrow aA' \mid \varepsilon & \end{array}$$

Пример множеств  $\text{FIRST}_1$  для нетерминалов грамматики  $G$ :

$$\begin{array}{ll} \text{FIRST}_1(S) = \{a\} & \text{FIRST}_1(B) = \{c, \varepsilon\} \\ \text{FIRST}_1(A) = \{a, \varepsilon\} & \text{FIRST}_1(S') = \{a, b, \varepsilon\} \\ \text{FIRST}_1(A') = \{a, b\} & \end{array}$$

Пример множеств FOLLOW<sub>1</sub> для нетерминалов грамматики  $G$ :

$$\begin{aligned}
 \text{FOLLOW}_1(S) &= \{\$ \} \\
 \text{FOLLOW}_1(S') &= \{\$ \} & (S \rightarrow aS') \\
 \text{FOLLOW}_1(A) &= \{b \} & (S' \rightarrow AbBS') \\
 \text{FOLLOW}_1(A') &= \{b \} & (A \rightarrow aA') \\
 \text{FOLLOW}_1(B) &= \{a, b, \$ \} & (S' \rightarrow AbBS', \varepsilon \in \text{FIRST}_1(S'))
 \end{aligned}$$

Управляющая таблица LL(k) анализатора заполняется следующим образом: продукции  $A \rightarrow \alpha$ ,  $\alpha \neq \varepsilon$  помещаются в ячейки  $(A, a)$ , где  $a \in \text{FIRST}_1(\alpha)$ , продукции  $A \rightarrow \alpha$  — в ячейки  $(A, a)$ , где  $a \in \text{FOLLOW}_1(A)$ , если  $\varepsilon \in \text{FIRST}_1(\alpha)$

**Пример 10.2.2.** Пример таблицы для грамматики  $S \rightarrow aSbS \mid \varepsilon$

| N | FIRST <sub>1</sub> | FOLLOW <sub>1</sub> | a                    | b                           | \$                          |
|---|--------------------|---------------------|----------------------|-----------------------------|-----------------------------|
| S | {a, ε}             | {b, \$}             | $S \rightarrow aSbS$ | $S \rightarrow \varepsilon$ | $S \rightarrow \varepsilon$ |

Однако, не для всех грамматик по множествам FIRST<sub>k</sub> и FOLLOW<sub>k</sub> возможно выбрать применяемую продукцию, а значит, нельзя однозначно построить таблицу, необходимую для работы алгоритма, поэтому данный алгоритм применим только для грамматик особого класса — LL(k).

**Определение 10.2.3.** LL(k) грамматика — грамматика, для которой на основании множеств FIRST<sub>k</sub> и FOLLOW<sub>k</sub> можно однозначно определить, какую продукцию применять.

Важно заметить, что при больших  $k$  управляющая таблица сильно разрастается, поэтому на практике данный алгоритм применим для небольших значений  $k$ .

Интерпретатор автомата принимает входную строку и построенную управляющую таблицу и работает следующим образом. В каждый момент времени конфигурация автомата это позиция во входной строке и стек. В начальный момент времени стек пуст, а позиция во входной строке соответствует её началу. На первом шаге в стек добавляются последовательно сперва символ конца строки, затем стартовый нетерминал. На каждом шаге анализируется существующая конфигурация и совершается одно из действий.

- Если текущая позиция — конец строки и вершина стека — символ конца строки, то успешно завершаем разбор.

- Если текущая вершина стека — терминал, то проверяем, что позиция в строке соответствует этому терминалу. Если да, то снимаем элемент со стека, сдвигаем позицию на единицу и продолжаем разбор. Иначе завершаем разбор с ошибкой.
- Если текущая вершина стека — нетерминал  $N_i$  и текущий входной символ  $t_j$ , то ищем в управляющей таблице ячейку с координатами  $(N_i, t_j)$  и записываем на стек содержимое этой ячейки.

**Пример 10.2.3.** Пример работы LL анализатора. Рассмотрим грамматику  $S \rightarrow aSbS \mid \varepsilon$  и выводимое слово  $\omega = abab$ .

Рассмотрим пошагово работу алгоритма, будем использовать таблицу, построенную в предыдущем примере:

1. Начало работы.

Стек: 

|    |
|----|
|    |
| \$ |

входное слово: 

|   |   |   |   |    |
|---|---|---|---|----|
| a | b | a | b | \$ |
|---|---|---|---|----|

Финальный символ лежит на стеке, а указатель указывает на первый символ слова.

2. кладем стартовый символ на стек

Стек: 

|    |
|----|
|    |
| S  |
| \$ |

входное слово: 

|   |   |   |   |    |
|---|---|---|---|----|
| a | b | a | b | \$ |
|---|---|---|---|----|

3. Ищем ячейку с координатами  $(S, a)$ , применяем продукцию из ячейки.

Стек: 

|    |
|----|
|    |
| a  |
| S  |
| b  |
| S  |
| \$ |

входное слово: 

|   |   |   |   |    |
|---|---|---|---|----|
| a | b | a | b | \$ |
|---|---|---|---|----|

4. Снимаем терминал  $a$  со стека и двигаем указатель.

Стек: 

|    |
|----|
|    |
| S  |
| b  |
| S  |
| \$ |

входное слово: 

|   |   |   |   |    |
|---|---|---|---|----|
| a | b | a | b | \$ |
|---|---|---|---|----|

5. Ищем ячейку с координатами  $(S, b)$ , применяем продукцию из ячейки.

|       |      |                |   |   |   |   |    |
|-------|------|----------------|---|---|---|---|----|
| Стек: |      | входное слово: | a | b | a | b | \$ |
|       | $b$  |                |   |   |   |   |    |
|       | $S$  |                |   |   |   |   |    |
|       | $\$$ |                |   |   |   |   |    |

6. Снимаем терминал  $b$  со стека и двигаем указатель.

|       |      |                |   |   |   |   |    |
|-------|------|----------------|---|---|---|---|----|
| Стек: |      | входное слово: | a | b | a | b | \$ |
|       | $S$  |                |   |   |   |   |    |
|       | $\$$ |                |   |   |   |   |    |

7. Ищем ячейку с координатами  $(S, a)$ , применяем продукцию из ячейки.

|       |      |                |   |   |   |   |    |
|-------|------|----------------|---|---|---|---|----|
| Стек: |      | входное слово: | a | b | a | b | \$ |
|       | $a$  |                |   |   |   |   |    |
|       | $S$  |                |   |   |   |   |    |
|       | $b$  |                |   |   |   |   |    |
|       | $S$  |                |   |   |   |   |    |
|       | $\$$ |                |   |   |   |   |    |

8. Снимаем терминал  $a$  со стека и двигаем указатель.

|       |      |                |   |   |   |   |    |
|-------|------|----------------|---|---|---|---|----|
| Стек: |      | входное слово: | a | b | a | b | \$ |
|       | $S$  |                |   |   |   |   |    |
|       | $b$  |                |   |   |   |   |    |
|       | $S$  |                |   |   |   |   |    |
|       | $\$$ |                |   |   |   |   |    |

9. Ищем ячейку с координатами  $(S, b)$ , применяем продукцию из ячейки.

|       |      |                |   |   |   |   |    |
|-------|------|----------------|---|---|---|---|----|
| Стек: |      | входное слово: | a | b | a | b | \$ |
|       | $b$  |                |   |   |   |   |    |
|       | $S$  |                |   |   |   |   |    |
|       | $\$$ |                |   |   |   |   |    |

10. Снимаем терминал  $b$  со стека и двигаем указатель.

|       |      |                |   |   |   |   |    |
|-------|------|----------------|---|---|---|---|----|
| Стек: |      | входное слово: | a | b | a | b | \$ |
|       | $S$  |                |   |   |   |   |    |
|       | $\$$ |                |   |   |   |   |    |

11. Ищем ячейку с координатами  $(S, \$)$ , применяем продукцию из ячейки.

|       |      |                |   |   |   |   |    |
|-------|------|----------------|---|---|---|---|----|
| Стек: |      | входное слово: | a | b | a | b | \$ |
|       | $\$$ |                |   |   |   |   |    |

12. Оказались в конце строки и на вершине стека символ конца — завершаем разбор.

Можно расширить данный алгоритм так, чтобы он строил дерево вывода. Дерево будет строиться сверху вниз, от корня к листьям. Для этого необходимо расширить шаги алгоритма.

- В ситуации, когда мы читаем очередной нетерминал (на вершине стека и во входе одинаковые терминалы), мы создаём лист с соответствующим терминалом.
- В ситуации, когда мы заменяем нетерминал на стеке на правую часть продукции в соответствии с управляющей таблицей, мы создаём нетерминальный узел соответствующий применяемой продукции.

Данное семейство всё так же не работает с леворекурсивными грамматиками и с неоднозначными грамматиками.

Таким образом, по некоторым грамматикам можно построить LL(k) анализатор (назовём их LL(k) грамматиками), но не по всем. С левой рекурсией, конечно, можно бороться, так как существуют алгоритмы устранения левой и скрытой левой рекурсии, а вот с неоднозначностями ничего не поделаешь.

## 10.3 Алгоритм Generalized LL

Можно построить анализатор, работающий с произвольными КС-грамматиками. Generalized LL (GLL) [56, 1]

Принцип работы остается абсолютно таким же как и для табличного LL:

- Сначала по грамматике строится *управляющая* таблица
- Затем построенная таблица команд и непосредственно анализируемое слово поступают на вход абстрактному интерпретатору.
- Для своей работы интерпретатор поддерживает некоторую вспомогательную структуру данных (стек для LL).
- Один шаг разбора состоит в том, чтобы рассмотреть текущую позицию в слове, применить все соответствующие ей правила из таблицы и при возможности сдвинуть позицию разбора вправо.



Где в этой схеме возникают ограничения на вид обрабатываемой грамматики для алгоритма LL? На самом первом шаге — при построении таблицы может возникнуть ситуация, когда одному нетерминалу  $N_j$  и последовательности  $first_k(N_j)$  соответствует несколько продукций грамматики. В этом случае грамматика признавалась не соответствующей классу LL(k) и отвергалась анализатором.

Теперь же мы разрешим такую ситуацию и в этом случае в ячейку таблицы будем записывать все продукции грамматики, соответствующие этой ячейке. Однако сразу же возникает вопрос — а что делать интерпретатору, когда при разборе ему необходимо применить правило, состоящее из нескольких продукций? Общий ответ такой — необходим некоторый вид недетерминизма, при котором интерпретатор мог бы “параллельно” обрабатывать несколько возможных вариантов синтаксического разбора.

Эти два свойства (модифицированная управляющая таблица и недетерминизм) суть главные принципиальные отличия GLL(k) от LL(k). Далее мы перейдем к рассмотрению непосредственно технической реализации описанного алгоритма.

Нам необходимо научиться задавать различные ветви (пути) синтаксического разбора и переключаться между ними. Заметим, что состояние любой ветви в любой момент времени суть следующее: необходимо распознать символ  $N_j \in N \cup \Sigma$  из продукции  $X$ , начиная с элемента слова под индексом  $i$ . Т.е. имеем позицию в слове и позицию символа в продукции. Последнее принято называть *слотом грамматики*.

**Определение 10.3.1.** Пусть  $G = \langle N, \Sigma, P, S \rangle$  — КС-грамматика. *Слотом грамматики  $G$*  (позицией грамматики  $G$ ) назовем пару из продукции  $X \in P$  и позиции  $0 \leq q \leq \text{length}(\text{body}(X))$  тела продукции  $X$ . При этом введем следующее обозначение  $X ::= \alpha \cdot \beta$ ,  $\alpha, \beta \in (N \cup \Sigma)^*$ , где  $\cdot$  указывает на позицию в продукции.

Описанная пара позиций уже однозначно задает состояние синтаксического разбора. Имеем множество состояний и переходов между ними — возникает естественное желание воспользоваться терминами графов для представления этой структуры. Такую конструкцию называют *граф-структурированный стек* или GSS (Graph Structured Stack), который впервые был предложен Масару Томитой [67] в контексте восходящего анализа. GSS будет являться рабочей структурой нашего нового интерпретатора вместо стека для LL. Состояние разбора вместе с узлом GSS мы будем называть *дескриптором*.

**Определение 10.3.2.** Пусть  $G = \langle N, \Sigma, P, S \rangle$  — КС-грамматика,  $X$  слот грамматики  $G$ ,  $i$  позиция в слове  $w$  над алфавитом  $\Sigma$ , а  $u$  узел GSS. *Дескриптором* назовём тройку  $(X, u, i)$ .

Есть несколько способов задания GSS для алгоритма GLL. Вариант, предложенный самими авторами алгоритма, оперирует непосредственно парами из позиции слова и слота грамматики в качестве состояний (и узлов графа) — такой метод является довольно простым и наглядным, но, как описано в работе [1], не самым эффективным. Предложим сразу чуть более оптимальное представление: заметим, что шаги разбора, соответствующие одному и тому же нетерминалу и позиции слова, должны выдавать один и тот же результат независимо от конкретной продукции грамматики, в которой стоит этот нетерминал. Поэтому заводить по узлу на каждый слот грамматики довольно избыточно — вместо этого в качестве состояния будет использовать пары из нетерминала и позиции слова, а позиции грамматики будем записывать на рёбрах.

Итак, мы научились задавать состояния с помощью дескрипторов, а также определились со вспомогательной структурой GSS. Теперь можно перейти к рассмотрению непосредственно самого алгоритма, суть которого довольно проста и напоминает BFS по неявному графу.

Дескриптор задает состояние, которое необходимо обработать. При этом мы без какой-либо дополнительной информации можем продолжить анализ входа из состояния, задаваемого этим дескриптором. В процессе обработки мы можем получить несколько новых состояний. Поэтому будем поддерживать множество  $R$  дескрипторов на обработку — на каждом шаге извлекаем один из множества, проводим анализ и кладем в множество новые полученные.

При каких условиях этот процесс будет конечен? Ну, например, если мы каждое состояние будем обрабатывать не более одного раза. И действительно, поскольку наш интерпретатор является “чистым” в том смысле, что для одного и того же состояния каждый раз будут получены одинаковые результаты, проводить анализ дважды не имеет смысла. Поэтому будем также поддерживать множество  $U$  всех полученных в ходе разбора дескрипторов, и добавлять в  $R$  только те, которых еще нет в  $U$ .

И наконец, заключительная и самая главная часть — как происходит обработка дескриптора? Пусть дескриптор имеет вид  $(X, u, i)$ , а входное слово обозначим  $W$ . Есть три возможных варианта, в зависимости от вида позиции грамматики  $X$  — разберем каждый из них по отдельности;

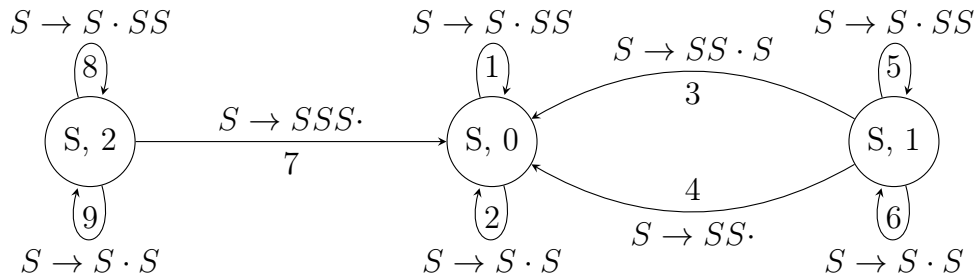
- $X ::= \alpha \cdot t\beta$ , т.е. указатель смотрит на терминал — в этом случае новых дескрипторов добавлено не будет. Если  $W[i] = t$ , то мы сдвигаем указатель слота, переходя к рассмотрению  $X ::= \alpha t \cdot \beta$ , и инкрементируем позицию  $i$  в слове. В противном же случае сразу переходим к следующему дескриптору, т.о. завершая текущую ветвь разбора.

- $X ::= \alpha \cdot A\beta$ , т.е. указатель смотрит на нетерминал. Нам нужен GSS узел  $v$  вида  $(A, i)$  и ребро  $(u, X ::= \alpha A \cdot \beta, v)$  (ребро из  $u$  в  $v$  с пометкой  $X ::= \alpha A \cdot \beta$ ). Если такой узел и ребро уже существуют в нашем GSS, берем их, иначе — создаём. Далее в  $R$  добавляем по дескриптору для узла  $v$  и каждого правила грамматики из ячейки управляющей таблицы для нетерминала  $A$  (конечно, если их еще не было в  $U$ ). На этом обработка текущего дескриптора завершается.
- $X ::= \alpha \cdot$ , т.е. указатель находится в конце продукции. Продукция разобрана, а значит, интерпретатору необходимо вернуться из разбора  $X$  к вызывающему правилу и продолжить разбор там (это, в некотором смысле, соответствует возврату из функции разбора нетерминала в методе рекурсивного спуска). По каждому исходящему ребру  $(u, Y, v)$  добавляем (если уже не существует) дескриптор  $(Y, v, i)$ .

Результатом синтаксического разбора является успех тогда и только тогда, когда был достигнут дескриптор вида  $(S ::= \alpha \cdot, s, n)$ , где слот грамматики представляет собой любое правило для аксиомы  $S$ , узел GSS  $s$  состоит из аксиомы  $S$  и 0, а позиция входного слова равна его длине  $n$ . Если же после разбора всех полученных дескрипторов указанный найден не был, результатом будет являться провал.

Давайте посмотрим, как такой алгоритм справится с неоднозначной грамматикой с леворекурсивным правилом.

**Пример 10.3.1.** Пусть грамматика  $G$  имеет вид  $S \rightarrow SSS \mid SS \mid a$ , а разбираемое слово  $w = aaa$ . Тогда GSS, соответствующий разбору  $S \Rightarrow SSS \Rightarrow aSS \Rightarrow aaS \Rightarrow aaa$ , будет выглядеть следующим образом (для удобства каждое ребро дополнительно пронумеровано):



Далее мы пошагово рассмотрим процесс его построения, а пока отметим несколько особенностей:

- Это *неполный* GSS. Для задачи синтаксического анализа такого достаточно, поскольку если в какой-то момент был достигнут финальный дескриптор, то обрабатывать все последующие уже не нужно. Однако, для задачи построения SPPF, как мы отметим далее, это уже не так, поскольку она требует агрегирования всех возможных путей разбора.
- Обратите особое внимание на наличие петель. Они как раз-таки и обеспечивают эффективную работу с леворекурсивными правилами, поскольку переиспользуются уже существующие узлы. При этом кратных петель, понятно, не создается, т.к. мы запоминаем все достигнутые дескрипторы в множестве  $U$  и дублирующих дескрипторов в рабочее множество  $R$  не добавляем.
- В GSS не создаются узлы, соответствующие разбору терминалов (например,  $a, 0$ ). В действительности так можно было бы сделать. Но тогда при обработке слота, указывающего на терминал, сначала бы создавался узел GSS, затем интерпретатор сверил бы терминал и символ в слове, после чего, если они совпали, произошел бы возврат из узла, а если нет, узел был бы отброшен и интерпретатор перешел бы к другому дескриптору. Таким образом, при любом случае сначала создается узел, затем выполняется проверка, после чего узел сразу отбрасывается. Для того, чтобы не создавать такие “одноразовые” узлы, проверка терминалов выполняется in-place.

Пронумеруем продукции и выпишем управляющую таблицу:

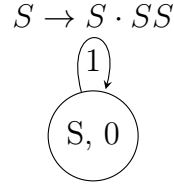
|                     |     |   |                    |       |    |
|---------------------|-----|---|--------------------|-------|----|
| $S \rightarrow SSS$ | (0) |   |                    |       |    |
| $S \rightarrow SS$  | (1) |   |                    |       |    |
| $S \rightarrow a$   | (2) |   |                    |       |    |
|                     |     | N | FIRST <sub>1</sub> | a     | \$ |
|                     |     | S | {a}                | 0,1,2 |    |

Разумеется, что конкретный порядок исполнения алгоритма будет зависеть, например, от используемой в качестве рабочего множества  $R$  структуры данных и от порядка обработки правил из ячейки управляющей таблицы. Рассмотрим лишь один из возможных вариантов:

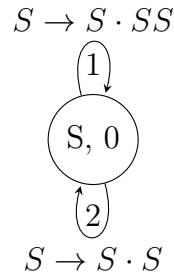
1. Для начала мы создаем узел GSS  $s_0 = (S, 0)$  и дескрипторы для правил из ячейки таблицы  $S, a$ :  $(S \rightarrow \cdot SSS, s_0, 0)$ ,  $(S \rightarrow \cdot SS, s_0, 0)$ ,  $(S \rightarrow \cdot a, s_0, 0)$ .



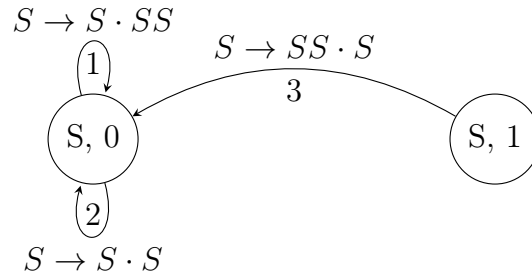
2. При обработке  $(S \rightarrow \cdot SSS, s_0, 0)$  образуются петля 1 и дескрипторы  $(S \rightarrow \cdot SSS, s_0, 0)$ ,  $(S \rightarrow \cdot SS, s_0, 0)$ ,  $(S \rightarrow \cdot a, s_0, 0)$ , которые уже содержатся в множестве  $U$  после шага 1 и поэтому не добавляются повторно.



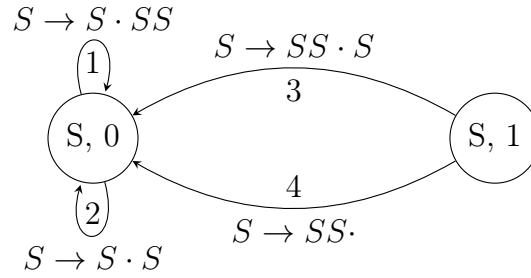
3. При обработке  $(S \rightarrow \cdot SS, s_0, 0)$  образуются петля 2, а в остальном аналогично шагу 2.



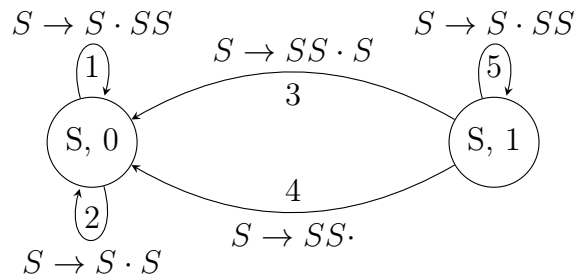
4. При обработке  $(S \rightarrow \cdot a, s_0, 0)$  мы распознаем терминал  $a$  на позиции 0 и, возвращаясь по петлям 1 и 2, добавляем дескрипторы  $(S \rightarrow S \cdot SS, s_0, 1)$ ,  $(S \rightarrow S \cdot S, s_0, 1)$ .
5. При обработке  $(S \rightarrow S \cdot SS, s_0, 1)$  образуем узел  $s_1 = (S, 1)$  с исходящим ребром 3 и добавляем дескрипторы  $(S \rightarrow \cdot SSS, s_1, 1)$ ,  $(S \rightarrow \cdot SS, s_1, 1)$ ,  $(S \rightarrow \cdot a, s_1, 1)$ .



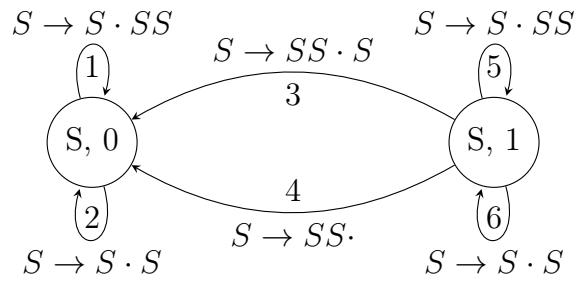
6. При обработке  $(S \rightarrow S \cdot S, s_0, 1)$  образуем ребро 4, новых дескрипторов не добавляется.



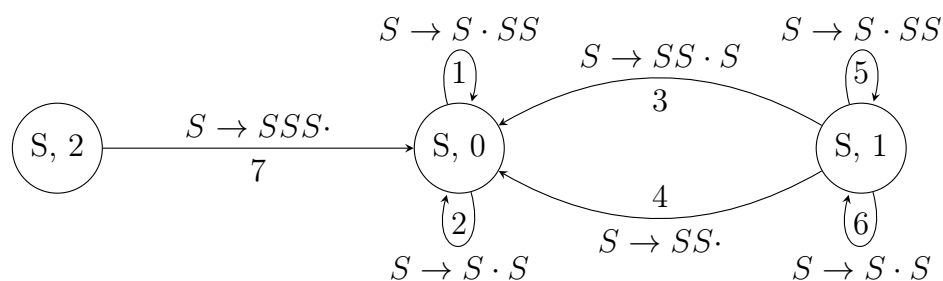
7. Обработка дескриптора  $(S \rightarrow \cdot SSS, s_1, 1)$  аналогична шагу 2 с добавлением петли 5.



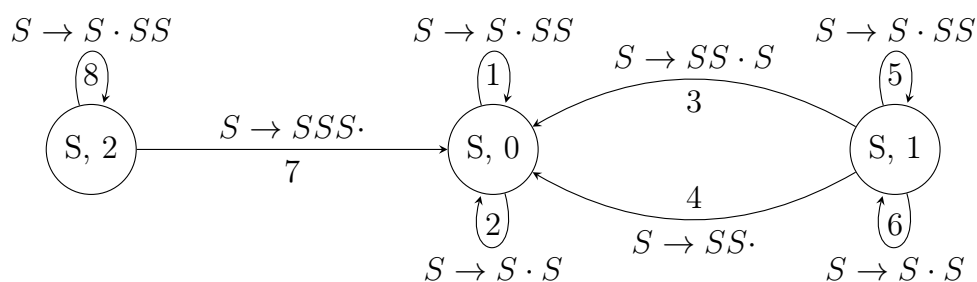
8. Обработка дескриптора  $(S \rightarrow \cdot SS, s_1, 1)$  аналогична шагу 3 с добавлением петли 6.



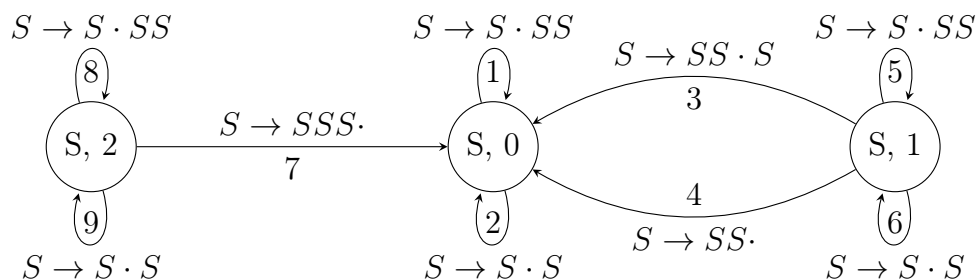
9. При обработке  $(S \rightarrow \cdot a, s_1, 1)$  мы распознаем терминал  $a$  на позиции 1 и, возвращаясь по ребрам 3 и 4, добавляем дескрипторы  $(S \rightarrow SS \cdot S, s_0, 2)$ ,  $(S \rightarrow SS \cdot, s_0, 2)$ , а также, возвращаясь по петлям 5 и 6, добавляем дескрипторы  $(S \rightarrow S \cdot SS, s_1, 2)$ ,  $(S \rightarrow S \cdot S, s_1, 2)$ .
10. При обработке  $(S \rightarrow SS \cdot S, s_0, 2)$  образуем узел  $s_2 = (S, 2)$  с исходящим ребром 7 и добавляем дескрипторы  $(S \rightarrow \cdot SSS, s_2, 2)$ ,  $(S \rightarrow \cdot SS, s_2, 2)$ ,  $(S \rightarrow \cdot a, s_2, 2)$ .



11. Обработка дескриптора  $(S \rightarrow \cdot SSS, s_2, 2)$  аналогична шагу 2 с добавлением петли 8.



12. Обработка дескриптора  $(S \rightarrow \cdot SS, s_2, 2)$  аналогична шагу 3 с добавлением петли 9.



13. При обработке  $(S \rightarrow \cdot a, s_2, 2)$  мы распознаем терминал  $a$  на позиции 2 и, возвращаясь по ребру 7, добавляем дескриптор  $(S \rightarrow SSS \cdot, s_0, 3)$ , а также, возвращаясь по петлям 8 и 9, добавляем дескрипторы  $(S \rightarrow S \cdot SS, s_2, 3)$ ,  $(S \rightarrow S \cdot S, s_2, 3)$ .
14. Мы достигли финального дескриптора  $(S \rightarrow SSS \cdot, s_0, 3)$ , синтаксический разбор успешен.

Внимательный читатель мог заметить, что если бы в этом примере шаг 4 был выполнен перед шагом 2, разбор довольно быстро бы завершился неудачей. Отсюда вытекает следующее наблюдение: если в какой-то момент из существующего узла появилось новое ребро, необходимо пересчитать все входящие в него пути.

Для построения SPPF требуется внести лишь несколько небольших добавлений:

1. В дескриптор необходимо добавить узел SPPF  $w$ , который будет представлять уже разобранный префикс.
2. Необходимо поддерживать множество  $P$  из элементов вида  $(u, z)$ , где  $u$  это узел GSS, а  $z$  соответствующий ему узел SPPF, для того, чтобы переиспользовать результаты разбора, ассоциированные с узлами GSS.
3. При обработке терминала  $t$  на позиции  $i$  ищется узел вида  $(t, i, i + 1)$ , либо создается, если такого еще нет.
4. При обработке нетерминала с помощью  $P$  ищется или при необходимости создается промежуточный узел вида  $(X, l, r)$ , где  $X$  соответствующий слот грамматики, а  $l$  и  $r$  узлы SPPF, отвечающие за разбор левой и правой частей слота соответственно.

Конкретные шаги построения SPPF будут зависеть от выбранного для него формата. Описание эффективного бинаризованного SPPF и детали его построения при выполнении GLL представлены в работе [1].

## 10.4 Алгоритм вычисления КС запросов на основе GLL

GLL довольно естественно обобщается на граф [29]: позициями входа теперь будем считать не индексы линейного слова, а вершины графа. В самом же алгоритме требуется внести лишь два небольших дополнения:

1. Теперь при обработке терминала “следующих” символов может быть несколько — рассматриваем каждый из них отдельно, сдвигаясь по соответствующему ребру. В результате, при одном чтении можем получить несколько новых дескрипторов, но они независимы, потому просто ставим их в рабочее множество  $R$ .



2. При обработке нетерминала, аналогично, правила управляющей таблицы применяются для каждого из “следующих” символов в графе. Соответственно новых дескрипторов будет сгенерировано больше, но все они по-прежнему независимы и просто добавляются в рабочее множество  $R$ .

Подробное описание алгоритма и псевдокод представлены в работе [29]. Существует ещё одно обобщение нисходящего синтаксического анализа для решения задачи КС достижимости [43], которое предполагает непосредственное обобщение LL(k) алгоритма, что приводит к аналогичному результату, однако теряется связь с некоторыми построениями.

Основанный на нисходящем анализе алгоритм поиска путей с контекстно-свободными ограничениями имеет следующие особенности.

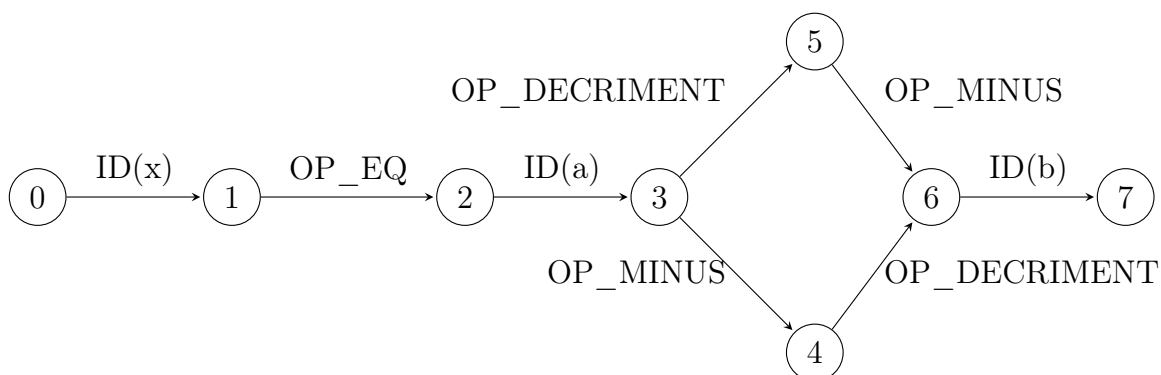
1. Необходимо явно задавать начальную вершину, поэтому хорошо подходит для задач поиска путей с одним источником (single-source) или с небольшим количеством источников. Для поиска путей между всеми парами вершин необходимо явным образом все указать стартовыми.
2. Является направленным сверху вниз — обходит граф последовательно начиная с указанной стартовой вершины и строит вывод, начиная со стартового нетерминала. Как следствие, в отличие от алгоритмов на основе линейной алгебры и Хеллингса, обойдёт только подграф, необходимый для построения ответа. В среднем это меньше, чем весь граф, который обрабатывается другими алгоритмами.
3. Естественным образом строит множество путей в виде сжатого леса разбора.
4. Использует существенно более тяжеловесные структуры данных и плохо распараллеливается (на практике). Как следствие, при решении задачи достижимости для всех пар путей проигрывает алгоритмам на основе линейной алгебры.

Частным случаем применения задачи КС достижимости является синтаксический анализ с неоднозначной токенизацией, то есть ситуацией, когда несколько пересекающихся подстрок во входной строке символов могут задавать разные лексические единицы и не возможно сделать однозначный выбор на этапе лексического анализа. Например, для строки  $x = a---b$  возможны несколько вариантов токенизации.

1. ID( $x$ ) OP\_EQ ID( $a$ ) OP\_MINUS OP\_DECRIMENT ID( $b$ )

2. ID(x) OP\_EQ ID(a) OP\_DECRIMENT OP\_MINUS ID(b)

В таком случае на вход синтаксическому анализатору можно подать DAG, содержащий все возможные варианты токенизации. Для нашего примера он может выглядеть следующим образом:



Далее будем проверять наличие пути из стартовой (нулевой) вершины в конечную (соответствующую концу строки). Если таких путей оказалось несколько, то нужны дополнительные средства для выбора нужного дерева разбора. Данная идея рассматривается в работе [57].

Напоследок сделаем небольшое замечание об эффективной реализации: в качестве рабочего множества  $R$  можно использовать несколько различных структур данных и, как правило, выбирают очередь. Однако иногда (в особенности для графов) лучше использовать стек дескрипторов, так как в этом случае выше локальность данных — мы кладём пачку дескрипторов, соответствующих исходящим рёбрам. И если граф представлен списком смежности, то исходящие будут храниться рядом и их лучше обработать сразу.

# Глава 11

## Алгоритм на основе восходящего анализа

В данном разделе будут рассмотрены алгоритмы восходящего синтаксического анализа LR-семейства, в том числе Generalized LR (GLR). Также будет рассмотрено обобщение алгоритма GLR для решения задачи поиска путей с контекстно-свободными ограничениями в графах.

### 11.1 Восходящий синтаксический анализ

Существует большое семейство  $LR(k)$  алгоритмов — алгоритм восходящего синтаксического анализа. Основная идея, лежащая в основе семейства, заключается в следующем: входная последовательность символов считывается слева направо с попутным добавлением в стек и выполнением сворачивания на стеке — замены последовательности терминалов и нетерминалов, лежащих наверху стека, на нетерминал, если существует соответствующее правило в исходной грамматике.

Как и в случае с LL используется магазинный автомат, управляемый таблицами, построенными по грамматике. При этом, у LR анализатора есть два типа команд:

1. `shift` — прочитать следующий символ входной последовательности, положив его в стек, и перейти в следующее состояние;
2. `reduce(k)` — применить  $k$ -ое правило грамматики, правая часть которого уже лежит на стеке: снимаем со стека правую часть продукции и кладем левую часть.

А управляющая таблица выглядит следующим образом.

| States | $t_0$ | ... | $t_a$ | ... | \$         | $N_0$ | ... | $N_b$ | ... |
|--------|-------|-----|-------|-----|------------|-------|-----|-------|-----|
| ...    | ...   | ... | ...   | ... | ...        | ...   | ... | ...   | ... |
| 10     | ...   | ... | $s_i$ | ... | $r_k$      | ...   | ... | $j$   | ... |
| ...    | ...   | ... | ...   | ... | <i>acc</i> | ...   | ... | ...   | ... |

Здесь

- $s_i$  — shift: перенести соответствующи символ в стек и перейти в состояние  $i$ .
- $r_k$  — reduce(k): в стеке накопилась правая часть продукции  $k$ , пора производить свёртку.
- $j$  — goto: выполняется после reduce. Сама по себе команда reduce не переводит автомат в новое состояние. Команда goto переведёт автомат в состояние  $j$ .
- *acc* — ассепт: разбор завершился успешно.

Если ячейка пустая и в процессе работы мы пропали в неё — значит произошла ошибка. Для детерминированной работы анализатора требуется, чтобы в каждой ячейке было не более одной команды. Если это не так, то говорят о возникновении конфликтов.

- shift-reduce — ситуация, когда не понятно, читать ли следующий символ или выполнить reduce. Например, если правая часть одного из правил является префиксом правой части другого правила:  $N \rightarrow w, M \rightarrow ww'$ .
- reduce-reduce — ситуация, когда не понятно, к какому правилу нужно применить reduce. Например, если есть два правила с одинаковыми правыми частями:  $N \rightarrow w, M \rightarrow w$ .

Принцип работы LR анализаторов следующий. Пусть у нас есть входная строка, LR-автомат со стеком и управляющая таблица. В начальный момент на стеке лежит стартовое состояние LR-автомата, позиция во входной строке соответствует её началу. На каждом шаге анализируется текущий символ входа и текущее состояние, в котором находится автомат, и совершается одно из действий:

- Если в управляющей таблице нет инструкции для текущего состояния автомата и текущего символа на входе, то завершаем разбор с ошибкой.
- Иначе выполняем одну из инструкций:
  - в случае *асс* — успешно завершаем разбор.
  - в случае *shift* — кладем на стек текущий символ входа, сдвигая при этом текущую позицию, и номер нового состояния. Переходим в новое состояние.
  - в случае *reduce(k)* — снимаем со стека  $2l$  элементов:  $l$  состояний и  $l$  терминалов/нетерминалов (где  $l$  — длина правой части  $k$ -ого правила), кладем на стек нетерминал левой части правила. Теперь на вершине стека у нас нетерминал  $N_a$ , а следующий элемент — состояние  $i$ . Если в ячейке  $(i, N_a)$  управляющей таблицы лежит состояние  $j$ , то кладем его на вершину стека. Иначе завершаем с ошибкой.

Разные алгоритмы из LR-семейства строят таблицы разными способами и, соответственно, могут избегать тех или иных конфликтов. Рассмотрим некоторых представителей.

### 11.1.1 LR(0) алгоритм

Данный алгоритм самый “слабый” из семейства — разбирает наименьший класс языков. Для построения используются LR(0) пункты.

**Определение 11.1.1.** LR(0) пункт (LR(0) item) — правило грамматики, в правой части которого имеется точка, отделяющая уже разобранный часть правила (слева от точки) от того, что еще предстоит распознать (справа от точки):  $A \rightarrow \alpha \cdot \beta$ , где  $A \rightarrow \alpha\beta$  — правило грамматики.  $\square$

Состояние LR(0) автомата — множество LR(0) пунктов. Для того чтобы их построить используется операция *closure* или *замыкание*.

**Определение 11.1.2.**  $\text{closure}(X) = \text{closure}(X \cup \{M \rightarrow \cdot \gamma \mid N_i \rightarrow \alpha \cdot M\beta \in X\})$   $\square$

**Определение 11.1.3.** Ядро — исходное множество пунктов, до применения к нему замыкания.  $\square$

Для перемещения точки в пункте используется функция *goto*.

**Определение 11.1.4.**  $\text{goto}(X, p) = \{N_j \rightarrow \alpha p \cdot \beta \mid N_j \rightarrow \alpha \cdot p\beta \in X\}$   $\square$

Теперь мы можем построить LR(0) автомат. Первым шагом необходимо расширить грамматику: добавить к исходной грамматике правило вида  $S' \rightarrow S\$$ , где  $S$  — стартовый нетерминал исходной грамматики,  $S'$  — новый стартовый нетерминал (не использовался ранее в грамматике),  $\$$  — маркер конца строки (не входил в терминальный алфавит исходной грамматики).

Далее строим автомат по следующим принципам.

- Состояния — множества пунктов.
- Переходы между состояниями осуществляются по символам грамматики.
- Начальное состояние —  $\text{closure}(\{S' \rightarrow \cdot S\})$ .
- Следующее состояние по текущему состоянию  $X$  и символу  $p$  вычисляются как  $\text{closure}(\text{goto}(X, p))$

Управляющая таблица по автомату строится следующим образом.

- $\text{acc}$  в ячейку, соответствующую финальному состоянию и  $\$$
- $s_i$  в ячейку  $(j, t)$ , если в автомате есть переход из состояния  $j$  по терминалу  $t$  в состояние  $i$
- $i$  в ячейку  $(j, N)$ , если в автомате есть переход из состояния  $j$  по нетерминалу  $N$  в состояние  $i$
- $r_k$  в ячейку  $(j, t)$ , если в состоянии  $j$  есть пункт  $A \rightarrow \alpha \cdot$ , где  $A \rightarrow \alpha$  —  $k$ -ое правило грамматики,  $t$  — терминал грамматики

### 11.1.2 SLR(1) алгоритм

SLR(1) анализатор отличается от LR(0) анализатора построением таблицы по автомату (автомат в точности как у LR(0)). А именно,  $r_k$  добавляется в ячейку  $(j, t)$ , если в состоянии  $j$  есть пункт  $A \rightarrow \alpha \cdot$ , где  $A \rightarrow \alpha$  —  $k$ -ое правило грамматики,  $t \in \text{FOLLOW}(A)$

### 11.1.3 CLR(1) алгоритм

Canonical LR(1), он же LR(1). Данный алгоритм является дальнейшим расширением SLR(1): к пунктам добавляются множества предпросмотра (lookahead).

**Определение 11.1.5.** Множество предпросмотра для правила  $P$  — терминалы, которые должны встретиться в выведенной строке сразу после строки, выводимой из данного правила.  $\square$

**Определение 11.1.6.** CLR пункт:  $[A \rightarrow \alpha \cdot \beta, \{t_0, \dots, t_n\}]$ , где  $t_0, \dots, t_n$  — множество предпросмотра для правила  $A \rightarrow \alpha\beta$ .  $\square$

**Определение 11.1.7.** Пусть дана грамматика  $G = \langle \Sigma, N, R, S \rangle$ .

$$\begin{aligned} closure(X) = closure(X \cup \{[B \rightarrow \cdot \delta, \{FIRST(\beta t_0), \dots, FIRST(\beta t_n)\}] \\ | B \rightarrow \beta \in R, [A \rightarrow \alpha \cdot B\beta, \{t_0, \dots, t_n\}] \in closure(X)\}) \end{aligned}$$

$\square$

Функция *goto* определяется аналогично LR(0), автомат строится по тем же принципам.

При построении управляющей таблицы усиливается правило добавления команды *redice*. А именно, добавляем  $r_k$  в ячейку  $(j, t_i)$ , если в состоянии  $j$  есть пункт  $[A \rightarrow \alpha \cdot, \{t_0, \dots, t_n\}]$ , где  $A \rightarrow \alpha$  —  $k$ -ое правило грамматики.

### 11.1.4 Примеры

Рассмотрим построение автоматов и таблиц для различных модификаций LR алгоритма.

Возьмем следующую грамматику:

$$\begin{array}{ll} 0) S & \rightarrow aSbS \\ 1) S & \rightarrow \varepsilon \end{array}$$

Расширим вышеупомянутую грамматику, добавив новый стартовый нетерминал  $S'$ , и далее будем работать с этой расширенной грамматикой:

$$\begin{array}{ll} 0) & S \rightarrow aSbS \\ 1) & S \rightarrow \varepsilon \\ 2) & S' \rightarrow S\$ \end{array}$$

**Пример 11.1.1.** Пример ядра и замыкания.

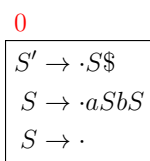
Возьмем правило 2 нашей грамматики, предположим, что мы только начинаем разбирать данное правило.

Ядром в таком случае является item исходного правила:  $S' \rightarrow \cdot S\$$

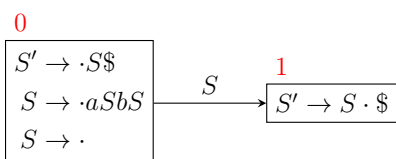
При замыкании добавятся ещё два item'а с правилами по выводу нетерминала 'S', поэтому получаем три item'а:  $S' \rightarrow \cdot S\$$ ,  $S \rightarrow \cdot aSbS$  и  $S \rightarrow \cdot \varepsilon$

**Пример 11.1.2.** Пример построения LR(0)-автомата для нашей грамматики с применением замыкания.

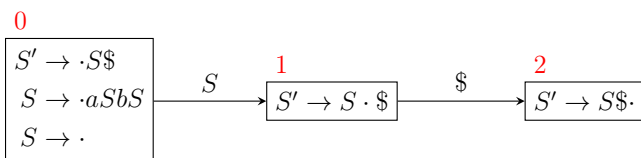
1. Добавляем стартовое состояние: item правила 0 и его замыкание (вместо item'а  $S \rightarrow \cdot \varepsilon$  будем писать  $S \rightarrow \cdot$ ).



2. По 'S' добавляем переход из стартового состояния в новое состояние 1.

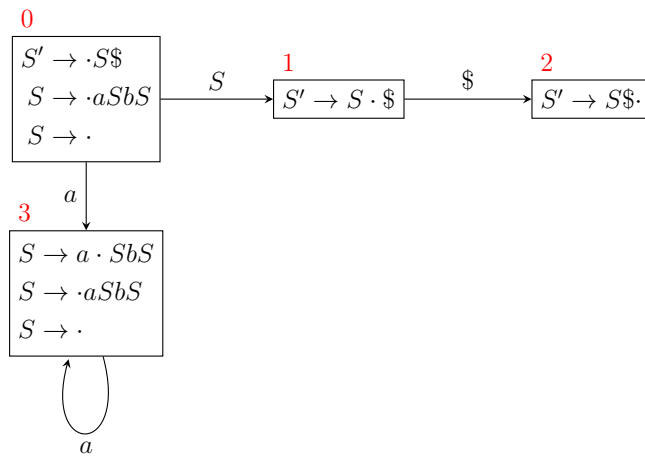


3. По '\$' добавляем переход из состояния 1 в новое состояние 2.

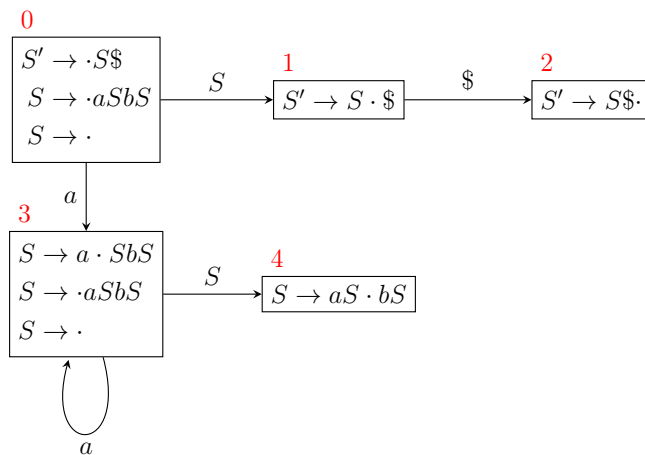


4. По 'a' добавляем переход из стартового состояния в новое состояние 3 и делаем его замыкание. Также добавляем переход по 'a' из этого состояния в себя же.

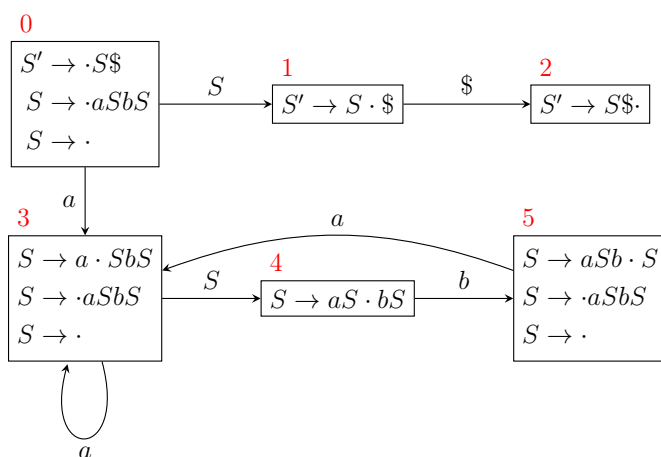




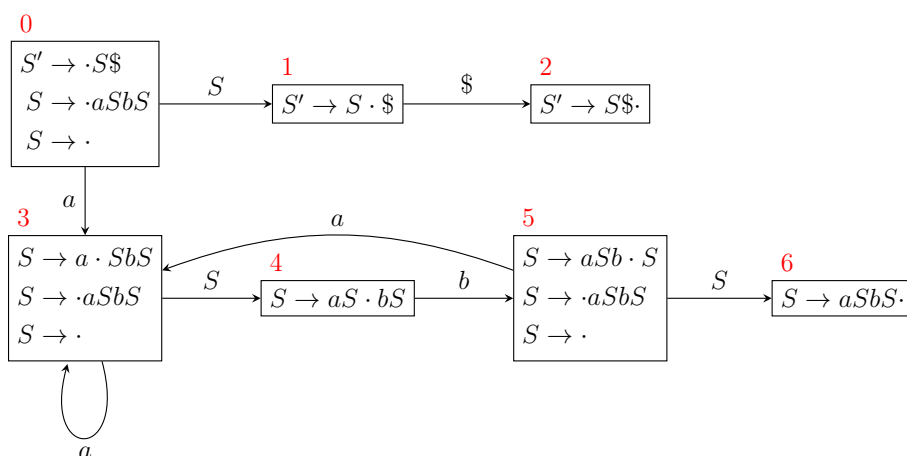
5. По 'S' добавляем переход из состояния 3 в новое состояние 4.



6. По 'b' добавляем переход из состояния 4 в новое состояние 5 и делаем его замыкание. Также добавляем переход по 'a' из этого состояния в состояние 3.



7. По 'S' добавляем переход из состояния 5 в новое состояние 6. Завершаем построение LR-автомата.



Далее будем использовать этот автомат для построения управляющей таблицы.

**Пример 11.1.3.** Пример управляющей LR(0) таблицы.

|   | a          | b     | \$    | S |
|---|------------|-------|-------|---|
| 0 | $s_3, r_1$ | $r_1$ | $r_1$ | 1 |
| 1 |            |       | acc   |   |
| 2 | $r_2$      | $r_2$ | $r_2$ |   |
| 3 | $s_3, r_1$ | $r_1$ | $r_1$ | 4 |
| 4 |            | $s_5$ |       |   |
| 5 | $s_3, r_1$ | $r_1$ | $r_1$ | 6 |
| 6 | $r_0$      | $r_0$ | $r_0$ |   |

Как видим, в данном случае в таблице присутствуют shift-reduce конфликты. В случае, когда не удаётся построить таблицу без конфликтов, говорят, что грамматика не LR(0).

□

**Пример 11.1.4.** Пример управляющей LR(1) таблицы. Автомат тот же, однако команды *reduce* расставляются с использованием FOLLOW.

$$FOLLOW_1(S) = \{b, \$\}$$

|   | a     | b     | \$    | S |
|---|-------|-------|-------|---|
| 0 | $s_3$ | $r_1$ | $r_1$ | 1 |
| 1 |       |       | acc   |   |
| 2 |       |       |       |   |
| 3 | $s_3$ | $r_1$ | $r_1$ | 4 |
| 4 |       | $s_5$ |       |   |
| 5 | $s_3$ | $r_1$ | $r_1$ | 6 |
| 6 |       | $r_0$ | $r_0$ |   |

В данном случае в таблице отсутствуют shift-reduce конфликты. То есть наша грамматика SLR(1), но не LR(0).

□

**Пример 11.1.5.** Пример LR-разбора входного слова abab\$ из языка нашей грамматики с использованием построенных ранее LR-автомата и управляющей таблицы.

1. Начало разбора. На стеке — стартовое состояние 0.

Вход: 

|   |   |   |   |    |
|---|---|---|---|----|
| a | b | a | b | \$ |
|---|---|---|---|----|

  
 Стек: 

|   |  |
|---|--|
| 0 |  |
|---|--|

2. Выполняем shift 3: сдвигаем указатель на входе, кладем на стек 'a', новое состояние 3 и переходим в него.

Вход: 

|   |   |   |   |    |
|---|---|---|---|----|
| a | b | a | b | \$ |
|---|---|---|---|----|

  
 Стек: 

|   |   |   |  |
|---|---|---|--|
| 0 | a | 3 |  |
|---|---|---|--|

3. Выполняем reduce 1 (кладем на стек 'S'), кладем новое состояние 4 и переходим в него.

|       |   |   |   |   |    |
|-------|---|---|---|---|----|
| Вход: | a | b | a | b | \$ |
| Стек: | 0 | a | 3 | S | 4  |

4. Выполняем shift 5: сдвигаем указатель на входе, кладем на стек 'b', новое состояние 5 и переходим в него.

|       |   |   |   |   |    |   |   |  |
|-------|---|---|---|---|----|---|---|--|
| Вход: | a | b | a | b | \$ |   |   |  |
| Стек: | 0 | a | 3 | S | 4  | b | 5 |  |

5. Выполняем shift 3.

|       |   |   |   |   |    |   |   |   |   |  |
|-------|---|---|---|---|----|---|---|---|---|--|
| Вход: | a | b | a | b | \$ |   |   |   |   |  |
| Стек: | 0 | a | 3 | S | 4  | b | 5 | a | 3 |  |

6. Выполняем reduce 1, кладем новое состояние 4 и переходим в него.

|       |   |   |   |   |    |   |   |   |   |   |   |  |
|-------|---|---|---|---|----|---|---|---|---|---|---|--|
| Вход: | a | b | a | b | \$ |   |   |   |   |   |   |  |
| Стек: | 0 | a | 3 | S | 4  | b | 5 | a | 3 | S | 4 |  |

7. Выполняем shift 5.

|       |   |   |   |   |    |   |   |   |   |   |   |   |   |
|-------|---|---|---|---|----|---|---|---|---|---|---|---|---|
| Вход: | a | b | a | b | \$ |   |   |   |   |   |   |   |   |
| Стек: | 0 | a | 3 | S | 4  | b | 5 | a | 3 | S | 4 | b | 5 |

8. Выполняем reduce 1, кладем новое состояние 6 и переходим в него.

|       |   |   |   |   |    |   |   |   |   |   |   |   |   |   |   |
|-------|---|---|---|---|----|---|---|---|---|---|---|---|---|---|---|
| Вход: | a | b | a | b | \$ |   |   |   |   |   |   |   |   |   |   |
| Стек: | 0 | a | 3 | S | 4  | b | 5 | a | 3 | S | 4 | b | 5 | S | 6 |

9. Выполняем reduce 0 (снимаем со стека 8 элементов и кладем 'S'), оказываемся в состоянии 5 и делаем переход в новое состояние 6 с добавлением его на стек.

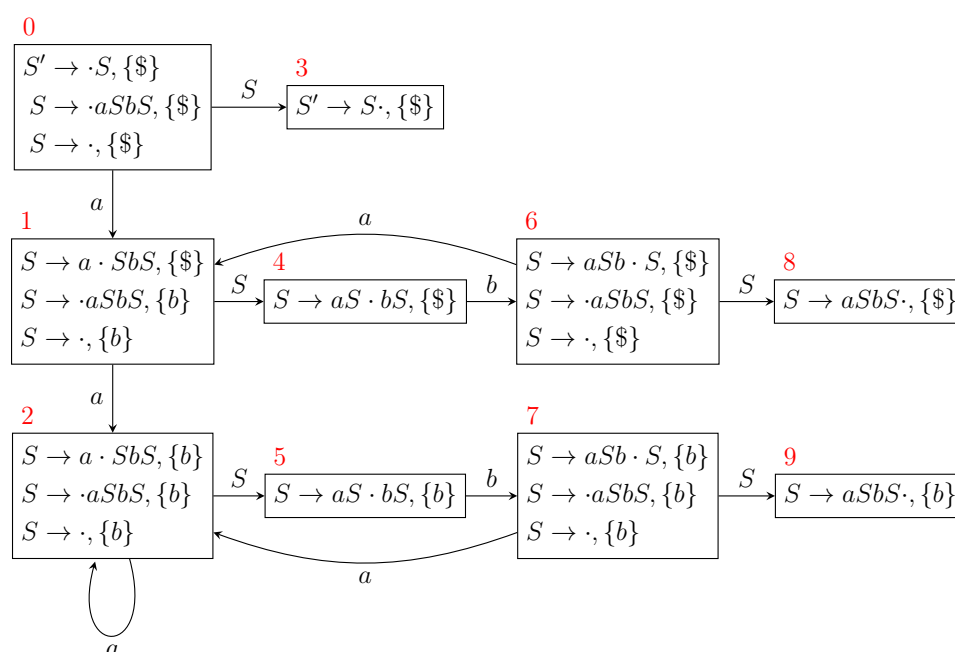
|       |   |   |   |   |    |   |   |   |   |
|-------|---|---|---|---|----|---|---|---|---|
| Вход: | a | b | a | b | \$ |   |   |   |   |
| Стек: | 0 | a | 3 | S | 4  | b | 5 | S | 6 |

10. Снова выполняем reduce 0, оказываемся в состоянии 0 и делаем переход в новое состояние 1 с добавлением его на стек. Заканчиваем разбор.

|       |   |   |   |   |    |  |  |  |  |
|-------|---|---|---|---|----|--|--|--|--|
| Вход: | a | b | a | b | \$ |  |  |  |  |
| Стек: | 0 | S | 1 |   |    |  |  |  |  |

□

**Пример 11.1.6.** Пример CLR автомата.



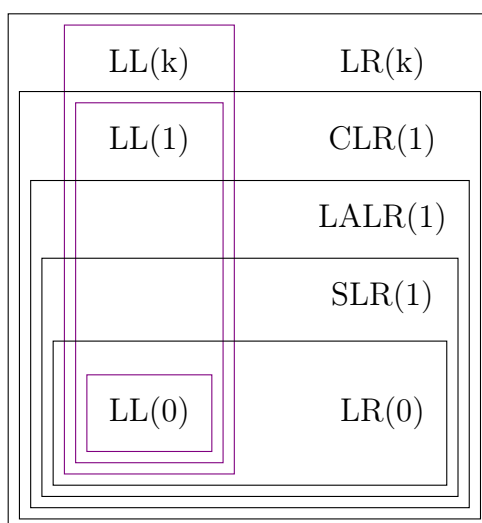
□

Существуют и другие модификации, например LALR(1)

На практике конфликты стараются решать ещё и на этапе генерации. Прикладные инструменты могут сгенерировать парсер по неоднозначной грамматике: из переноса или свёртки выбирать перенос, из нескольких свёрток — первую в каком-то порядке (обычно в порядке появления соответствующих productions в грамматике).

### 11.1.5 Сравнение классов LL и LR

Иерархию языков, распознаваемых различными классами алгоритмов, можно представить следующим образом.



Из диаграммы видно, что класс языков, распознаваемых  $LL(k)$  алгоритмом уже, чем класс языков, распознаваемый  $LR(k)$  алгоритмом, при любом конечном  $k$ . Приведём несколько примеров.

1.  $L = \{a^m b^n c \mid m \geq n \geq 0\}$  является  $LR(0)$ , но для него не существует  $LL(1)$  грамматики.
2.  $L = \{a^n b^n + a^n c^n \mid n > 0\}$  является  $LR$ , но не  $LL$ .
3. Больше примеров можно найти в работе Джона Битти [13].

## 11.2 GLR и его применение для КС запросов

Алгоритм  $LR$  довольно эффективен, однако позволяет работать не со всеми КС-грамматиками, а только с их подмножеством  $LR(k)$ . Если грамматика находится за рамками допускаемого класса, некоторые ячейки управляющей таблицы могут содержать несколько значений. В этом случае грамматика отвергалась анализатором.

Чтобы допустить множественные значения в ячейках управляющей таблицы, потребуется некоторый вид недетерминизма, который даст возможность

анализатору обрабатывать несколько возможных вариантов синтаксического разбора параллельно. Именно это и предлагает анализатор Generalized LR (GLR) [66]. Далее мы рассмотрим общий принцип работы, проиллюстрируем его с помощью примера, а также рассмотрим модификации GLR.

### 11.2.1 Классический GLR алгоритм

Впервые GLR парсер был представлен Масару Томитой в 1987 [66]. В целом, алгоритм работы идентичен LR той разницей, что управляющая таблица модифицирована таким образом, чтобы допускать множественные значения в ячейках. Интерпретатор автомата изменён соответствующим образом.

Для того, чтобы избежать дублирования информации при обработке неоднозначностей, стоит использовать более сложную структуру стека: *граф-структурированный стек* или (*GSS*, Graph Structured Stack). Это направленный граф, в котором вершины соответствуют элементам стека, а ребра их соединяют по правилам управляющей таблицы. У каждой вершины может быть несколько входящих и исходящих дуг: таким образом реализуется то объединение одинаковых состояний и ветвление в случае неоднозначности.

**Пример 11.2.1.** Рассмотрим пример GLR разбора с использованием GSS.

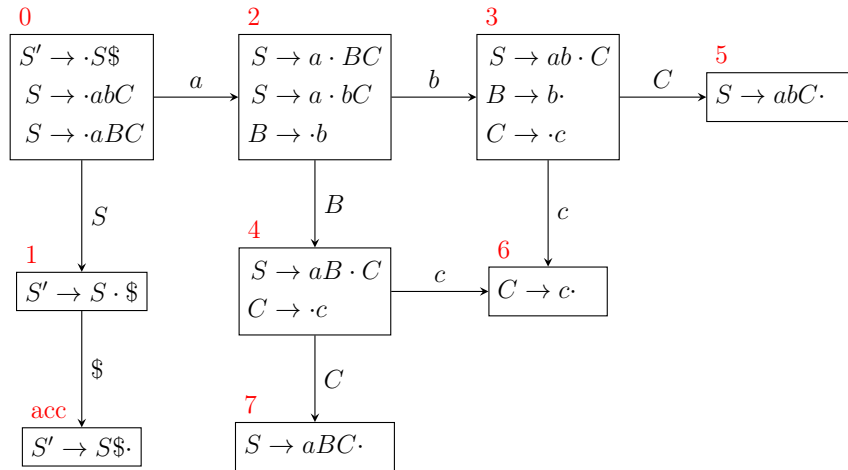
Возьмем грамматику  $G$  следующего вида:

0.  $S' \rightarrow S\$$
1.  $S \rightarrow abC$
2.  $S \rightarrow aBC$
3.  $B \rightarrow b$
4.  $C \rightarrow c$

Входное слово  $w$ :

$$w = abc\$$$

Построим для данной грамматики LR автомат:



И управляющую таблицу:

|   | a     | b     | c          | \$    | B | C | S |
|---|-------|-------|------------|-------|---|---|---|
| 0 | $s_2$ |       |            |       |   | 1 |   |
| 1 |       |       |            | acc   |   |   |   |
| 2 |       | $s_3$ |            |       | 4 |   |   |
| 3 |       |       | $s_6, r_3$ |       |   | 5 |   |
| 4 |       |       | $s_6$      |       |   | 7 |   |
| 5 |       |       |            | $r_1$ |   |   |   |
| 6 |       |       |            | $r_4$ |   |   |   |
| 7 |       |       |            | $r_2$ |   |   |   |

Разберем слово  $w$  с помощью алгоритма GLR. Использована следующая аннотация: вершины-состояния обозначены кругами, вершины-символы — прямоугольниками.

1. Инициализируем GSS стартовым состоянием  $v_0$ :

Вход: 

|   |   |   |    |
|---|---|---|----|
| a | b | c | \$ |
|---|---|---|----|

 GSS:  $\overset{v_0}{\textcircled{0}}$

2. Видим входной символ 'a', ищем соответствующую ему операцию в управляющей таблице — *shift* 2, строим новый узел  $v_1$ :

Вход: 

|   |   |   |    |
|---|---|---|----|
| a | b | c | \$ |
|---|---|---|----|

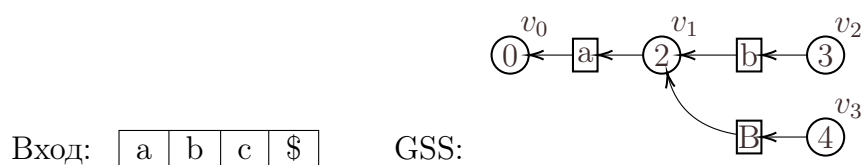
 GSS:  $\overset{v_0}{\textcircled{0}} \leftarrow \text{a} \leftarrow \overset{v_1}{\textcircled{2}}$



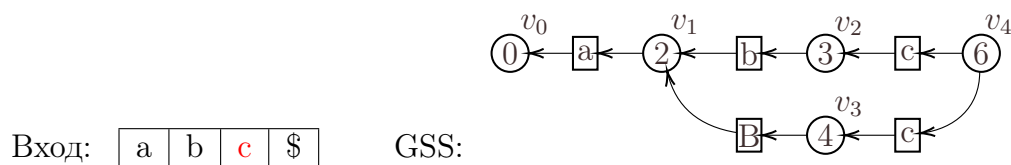
3. Повторяем для символа 'b', операции *shift* 3 и узла  $v_2$ :



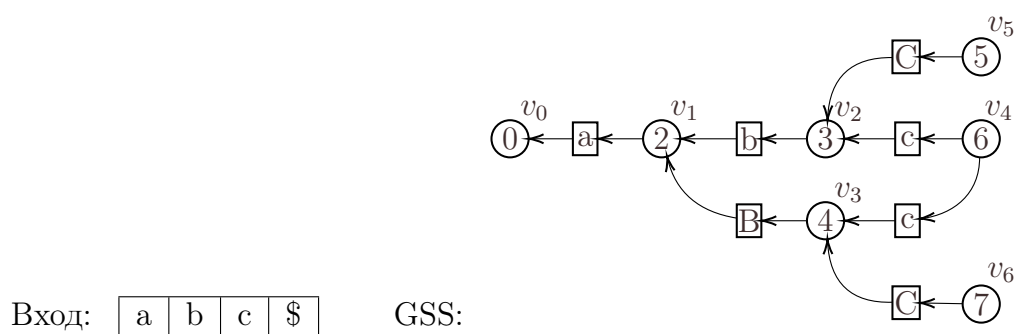
4. При обработке узла  $v_3$  у нас возникает конфликт shift-reduce:  $s_6$ ,  $r_3$ . Мы смотрим на вершины, смежные  $v_2$ , на управляющую таблицу и на правило вывода под номером 3 для поиска альтернативного построения стека. Находим *goto* 4 и строим вершину  $v_3$  с соответствующим переходом по нетерминалу  $B$  из  $v_1$  (т.к. количество символов в правой части правила вывода 3 равняется 1, значит мы в дереве опустимся на глубину 1 по вершинам-состояниям):



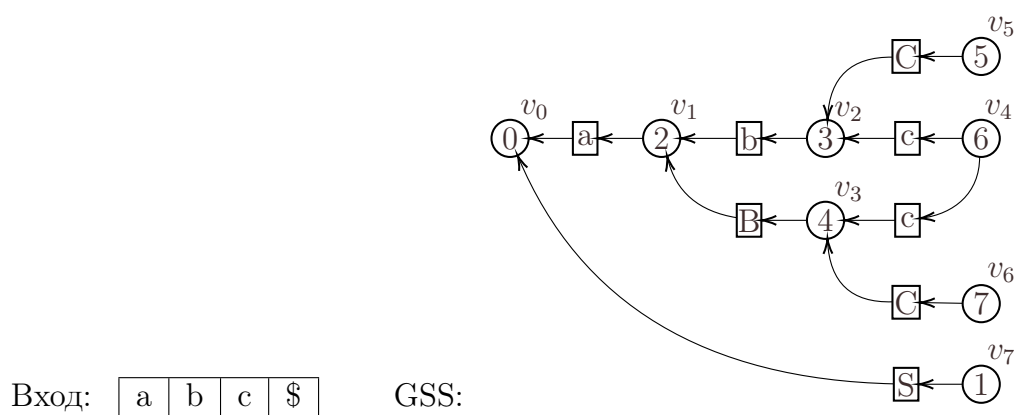
5. Читаем символ 'c' и ищем в управляющей таблице переходы из состояний 3 и 4 (так как узлы  $v_2$  и  $v_3$  находятся на одном уровне, то есть были построены после чтения одного символа из входного слова). Таким переходом оказывается  $s_6$  в обоих случаях, поэтому соединяем узел  $v_4$  с обоими рассмотренными узлами:



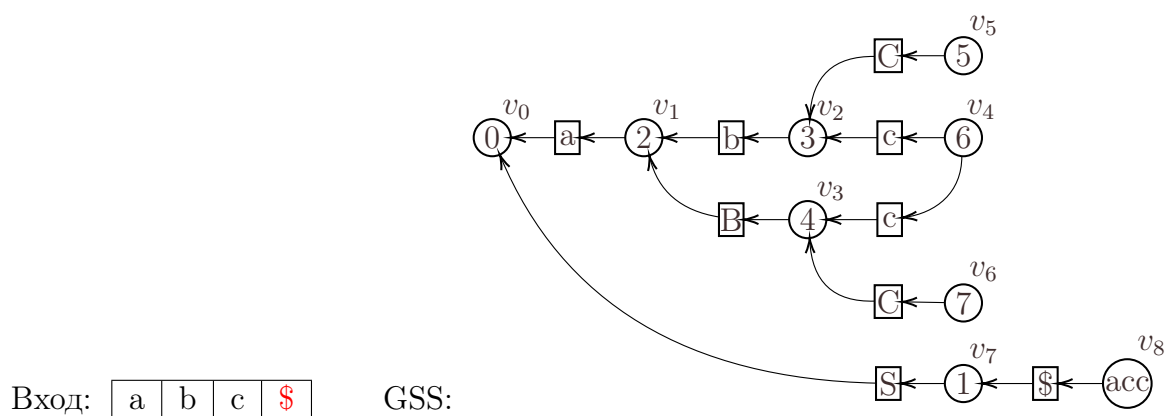
6. При обработке узла  $v_4$  находим соответствующую 6-ому состоянию редукцию по правилу 4. Его правая часть содержит один символ 'c', 2 вершины-символа с которым достижимы из  $v_4$ . Находим вершины-состояния, которые смежны с этими вершинами-символами и обрабатываем переходы по левой части правила 4. Такими переходами по нетерминалу  $C$  оказываются 5 и 7. Строим соответствующие им вершины  $v_5$  и  $v_6$ :



7. При обработке узлов  $v_5$  и  $v_6$  находим редукции с символом 'S' в левой части и тремя символами в правой. Возвращаемся на 3 вершины-состояния назад и строим вершину  $v_7$  с переходом по  $S$ :



8. Наконец, обрабатывая вершину  $v_7$ , читаем символ '\$' и строим узел  $v_8$ , который соответствует допускаям состоянием:



### 11.2.2 Модификации GLR

Алгоритм, представленный Томитой имел большой недостаток: он корректно работал не со всеми КС грамматиками, хоть и расширял класс допустимых LR анализаторами. Объем потребляемой памяти классическим GLR можно оценить как  $O(n^3)$  с учетом оптимизаций, о которых говорилось ранее.

Спустя некоторое время после публикации Томита-парсера, Элизабет Скотт и Эндрю Джонстоун представили *RNGLR* (Right Nulled GLR) [55] — модифицированная версия GLR, которая решала проблему скрытых рекурсий. Это позволило расширить класс допускаемых грамматик до КС. Однако объем потребляемой памяти можно оценить сверху уже полиномом  $O(n^{k+1})$ , где  $k$  — длина самого длинного правила грамматики, что несколько ухудшило оценку классического GLR.

С этой проблемой справился BRNGLR (Binary RNGLR) [58]. За счет бинаризации удалось получить кубическую оценку сложности и при этом также, как и RNGLR, допускать все КС грамматики.

Кроме того, GLR довольно естественно обобщается до решения задачи поиска путей с КС ограничениями. Это происходит следующим образом: элементами во входной структуре теперь будем считать не позиции символа в слове, а вершины графа (то есть “позиции” и множество смежных вершин). Это приводит к тому, что при применении операции shift, следующих символов может быть несколько и каждый из них должен быть рассмотрен отдельно, сдвигаясь по соответствующему ребру и проходя входной граф в ширину. Подробное описание алгоритма и псевдокод представлены в работе [69].



# Литература

- [1] A. Afroozeh and A. Izmaylova. Faster, practical gll parsing. In B. Franke, editor, *Compiler Construction*, pages 89–108, Berlin, Heidelberg, 2015. Springer Berlin Heidelberg.
- [2] M. R. Albrecht, G. V. Bard, and W. Hart. Efficient multiplication of dense matrices over  $\text{GF}(2)$ . *CoRR*, abs/0811.1714, 2008.
- [3] J. Alman and V. V. Williams. A refined laser method and faster matrix multiplication, 2020.
- [4] J. W. Anderson, P. Tataru, J. Staines, J. Hein, and R. Lyngsø. Evolving stochastic context-free grammars for RNA secondary structure prediction. *BMC Bioinformatics*, 13(1):78, 2012.
- [5] R. Axelsson and M. Lange. Formal language constrained reachability and model checking propositional dynamic logics. In *International Workshop on Reachability Problems*, pages 45–57. Springer, 2011.
- [6] R. Azimov and S. Grigorev. Context-free path querying by matrix multiplication. In *Proceedings of the 1st ACM SIGMOD Joint International Workshop on Graph Data Management Experiences & Systems (GRADES) and Network Data Analytics (NDA)*, GRADES-NDA '18, pages 5:1–5:10, New York, NY, USA, 2018. ACM.
- [7] R. Azimov and S. Grigorev. Path querying using conjunctive grammars. *Proceedings of the Institute for System Programming of the RAS*, 30:149–166, 01 2018.
- [8] N. Bansal and R. Williams. Regularity lemmas and combinatorial algorithms. In *2009 50th Annual IEEE Symposium on Foundations of Computer Science*, pages 745–754, 2009.
- [9] J. Baras and G. Theodorakopoulos. Path problems in networks. In *Path Problems in Networks*, 2010.

- [10] C. Barrett, K. Bisset, M. Holzer, G. Konjevod, M. Marathe, and D. Wagner. Label constrained shortest path algorithms: An experimental evaluation using transportation networks. *March*, 9:2007, 2007.
- [11] C. Barrett, R. Jacob, and M. Marathe. Formal-language-constrained path problems. *SIAM Journal on Computing*, 30(3):809–837, 2000.
- [12] O. Bastani, S. Anand, and A. Aiken. Specification inference using context-free language reachability. In *ACM SIGPLAN Notices*, volume 50, pages 553–566. ACM, 2015.
- [13] J. C. Beatty. Two iteration theorems for the  $ll(k)$  languages. *Theoretical Computer Science*, 12(2):193 – 228, 1980.
- [14] S. Billot and B. Lang. The structure of shared forests in ambiguous parsing. In *Proceedings of the 27th Annual Meeting on Association for Computational Linguistics*, ACL '89, pages 143–151, Stroudsburg, PA, USA, 1989. Association for Computational Linguistics.
- [15] D. A. Bini, M. Capovani, F. Romani, and G. Lotti.  $o(n^{2.7799})$  complexity for approximate matrix multiplication. *Information Processing Letters*, 8(5):234–235, 1979.
- [16] P. G. Bradford. Efficient exact paths for dyck and semi-dyck labeled path reachability (extended abstract). In *2017 IEEE 8th Annual Ubiquitous Computing, Electronics and Mobile Communication Conference (UEMCON)*, pages 247–253, Oct 2017.
- [17] T. M. Chan. All-pairs shortest paths with real weights in  $o(n^3/\log n)$  time. *Algorithmica*, 50(2):236–243, Feb 2008.
- [18] T. M. Chan. More algorithms for all-pairs shortest paths in weighted graphs. *SIAM Journal on Computing*, 39(5):2075–2089, 2010.
- [19] D. Coppersmith and S. Winograd. On the asymptotic complexity of matrix multiplication. *SIAM Journal on Computing*, 11(3):472–492, 1982.
- [20] D. Coppersmith and S. Winograd. Matrix multiplication via arithmetic progressions. *Journal of symbolic computation*, 9(3):251–280, 1990.
- [21] S. Dalton, N. Bell, L. Olson, and M. Garland. Cusp: Generic parallel algorithms for sparse matrix and graph computations, 2014. Version 0.5.0.
- [22] D. Das, M. Koucký, and M. Saks. Lower bounds for combinatorial algorithms for boolean matrix multiplication, 2018.

- [23] T. A. Davis. Algorithm 9xx: Suitesparse:graphblas: graph algorithms in the language of sparse linear algebra. 2018.
- [24] W. Dobosiewicz. A more efficient algorithm for the min-plus multiplication. *International journal of computer mathematics*, 32(1-2):49–60, 1990.
- [25] W. Dyrka, M. Pyzik, F. Coste, and H. Talibart. Estimating probabilistic context-free grammars for proteins using contact map constraints. *PeerJ*, 7:e6559, Mar. 2019.
- [26] N. El abbadi. An efficient storage format for large sparse matrices based on quadtree. *International Journal of Computer Applications*, 105:25–30, 11 2014.
- [27] R. W. Floyd. Algorithm 97: shortest path. *Communications of the ACM*, 5(6):345, 1962.
- [28] M. L. Fredman. New bounds on the complexity of the shortest path problem. *SIAM Journal on Computing*, 5(1):83–89, 1976.
- [29] S. Grigorev and A. Ragozina. Context-free path querying with structural representation of result. In *Proceedings of the 13th Central & Eastern European Software Engineering Conference in Russia, CEE-SECR '17*, pages 10:1–10:7, New York, NY, USA, 2017. ACM.
- [30] Y. Han. Improved algorithm for all pairs shortest paths. *Information Processing Letters*, 91(5):245–250, 2004.
- [31] J. Hellings. Conjunctive context-free path queries. In *Proceedings of ICDT'14*, pages 119–130, 2014.
- [32] J. Hellings. Path results for context-free grammar queries on graphs. *ArXiv*, abs/1502.02242, 2015.
- [33] K. Hemerik. Towards a taxonomy for ecfg and rrpq parsing. *Proceedings of the 3rd International Conference on Language and Automata Theory and Applications*, pages 410–421, 2009.
- [34] M. Holzer, M. Kutrib, and U. Leiter. Nodes connected by path languages. In G. Mauri and A. Leporati, editors, *Developments in Language Theory*, pages 276–287, Berlin, Heidelberg, 2011. Springer Berlin Heidelberg.
- [35] J. E. Hopcroft and J. D. Ullman. *Introduction to Automata Theory, Languages, and Computation*. Addison-Wesley Publishing Company, 1979.

- [36] A. K. Joshi and Y. Schabes. *Tree-Adjoining Grammars*, pages 69–123. Springer Berlin Heidelberg, Berlin, Heidelberg, 1997.
- [37] J. Kasai, B. Frank, T. McCoy, O. Rambow, and A. Nasr. TAG parsing with neural networks and vector representations of supertags. In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, pages 1712–1722, Copenhagen, Denmark, Sept. 2017. Association for Computational Linguistics.
- [38] J. Kasai, R. Frank, P. Xu, W. Merrill, and O. Rambow. End-to-end graph-based TAG parsing with neural networks. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*, pages 1181–1194, New Orleans, Louisiana, June 2018. Association for Computational Linguistics.
- [39] J. Kepner and J. Gilbert. *Graph Algorithms in the Language of Linear Algebra*. Society for Industrial and Applied Mathematics, 2011.
- [40] O. Kupferman and G. Vardi. Eulerian paths with regular constraints. In *41st International Symposium on Mathematical Foundations of Computer Science (MFCS 2016)*. Schloss Dagstuhl-Leibniz-Zentrum fuer Informatik, 2016.
- [41] M. J. Kusner, B. Paige, and J. M. Hernández-Lobato. Grammar variational autoencoder. In *Proceedings of the 34th International Conference on Machine Learning - Volume 70, ICML’17*, page 1945–1954. JMLR.org, 2017.
- [42] L. Lee. Fast context-free grammar parsing requires fast boolean matrix multiplication. *J. ACM*, 49(1):1–15, Jan. 2002.
- [43] C. M. Medeiros, M. A. Musicante, and U. S. Costa. Ll-based query answering over rdf databases. *Journal of Computer Languages*, 51:75 – 87, 2019.
- [44] N. Mishin, I. Sokolov, E. Spirin, V. Kutuev, E. Nemchinov, S. Gorbatyuk, and S. Grigorev. Evaluation of the context-free path querying algorithm based on matrix multiplication. In *Proceedings of the 2Nd Joint International Workshop on Graph Data Management Experiences & Systems (GRADES) and Network Data Analytics (NDA)*, GRADES-NDA’19, pages 12:1–12:5, New York, NY, USA, 2019. ACM.
- [45] A. Okhotin. Conjunctive grammars. *Journal of Automata, Languages and Combinatorics*, 6(4):519–535, 2001.



- [46] A. Okhotin. On the closure properties of linear conjunctive languages. *Theor. Comput. Sci.*, 299(1-3):663–685, 2003.
- [47] V. Y. Pan. Strassen’s algorithm is not optimal trilinear technique of aggregating, uniting and canceling for constructing fast algorithms for matrix operations. In *19th Annual Symposium on Foundations of Computer Science (sfcs 1978)*, pages 166–176. IEEE, 1978.
- [48] P. Pratikakis, J. S. Foster, and M. Hicks. Existential label flow inference via cfl reachability. In *SAS*, volume 6, pages 88–106. Springer, 2006.
- [49] J. Rehof and M. Fähndrich. Type-base flow analysis: From polymorphic subtyping to cfl-reachability. *SIGPLAN Not.*, 36(3):54–66, Jan. 2001.
- [50] J. G. Rekers. *Parser generation for interactive environments*. PhD thesis, Universiteit van Amsterdam, 1992.
- [51] T. Reps. Program analysis via graph reachability. In *Proceedings of the 1997 International Symposium on Logic Programming, ILPS ’97*, pages 5–19, Cambridge, MA, USA, 1997. MIT Press.
- [52] T. Reps, S. Horwitz, M. Sagiv, and G. Rosay. Speeding up slicing. In *Proceedings of the 2nd ACM SIGSOFT Symposium on Foundations of Software Engineering, SIGSOFT ’94*, pages 11–20, New York, NY, USA, 1994. Association for Computing Machinery.
- [53] B. Roy. Transitivité et connexité. *Comptes Rendus Hebdomadaires Des Seances De L Academie Des Sciences*, 249(2):216–218, 1959.
- [54] A. Schönhage. Partial and total matrix multiplication. *SIAM Journal on Computing*, 10(3):434–455, 1981.
- [55] E. Scott and A. Johnstone. Right nulled glr parsers. *ACM Trans. Program. Lang. Syst.*, 28(4):577–618, July 2006.
- [56] E. Scott and A. Johnstone. Gll parsing. *Electron. Notes Theor. Comput. Sci.*, 253(7):177–189, Sept. 2010.
- [57] E. Scott and A. Johnstone. Multiple lexicalisation (a java based study). In *Proceedings of the 12th ACM SIGPLAN International Conference on Software Language Engineering, SLE 2019*, page 71–82, New York, NY, USA, 2019. Association for Computing Machinery.
- [58] E. Scott, A. Johnstone, and R. Economopoulos. Brnglr: A cubic tomita-style glr parsing algorithm. *Acta Inf.*, 44(6):427–461, Sept. 2007.

- [59] H. Seki, T. Matsumura, M. Fujii, and T. Kasami. On multiple context-free grammars. *Theoretical Computer Science*, 88(2):191 – 229, 1991.
- [60] P. Sevon and L. Eronen. Subgraph queries by context-free grammars. *Journal of Integrative Bioinformatics*, 5, 06 2008.
- [61] E. Shemetova and S. Grigorev. Path querying on acyclic graphs using boolean grammars. *Proceedings of the Institute for System Programming of RAS*, 31(4):211–226, 2019.
- [62] V. Strassen. Gaussian elimination is not optimal. *Numerische mathematik*, 13(4):354–356, 1969.
- [63] T. Takaoka. A new upper bound on the complexity of the all pairs shortest path problem. *Information Processing Letters*, 43(4):195–199, 1992.
- [64] T. Takaoka. A faster algorithm for the all-pairs shortest path problem and its application. In *Computing and Combinatorics*, volume 3106, pages 278–289, 2004.
- [65] T. Takaoka. An  $o(n^3 \log \log n / \log n)$  time algorithm for the all-pairs shortest path problem. *Information Processing Letters*, 96:155–161, 2005.
- [66] M. Tomita. An efficient augmented-context-free parsing algorithm. *Computational Linguistics*, 13:31–46, 1987.
- [67] M. Tomita. Graph-structured stack and natural language parsing. In *26th Annual Meeting of the Association for Computational Linguistics*, pages 249–257, Buffalo, New York, USA, June 1988. Association for Computational Linguistics.
- [68] L. G. Valiant. General context-free recognition in less than cubic time. *J. Comput. Syst. Sci.*, 10(2):308–315, Apr. 1975.
- [69] E. Verbitskaia, S. Grigorev, and D. Avdyukhin. Relaxed parsing of regular approximations of string-embedded languages. In M. Mazzara and A. Voronkov, editors, *Perspectives of System Informatics*, pages 291–302, Cham, 2016. Springer International Publishing.
- [70] C. B. Ward and N. M. Wiegand. Complexity results on labeled shortest path problems from wireless routing metrics. *Comput. Netw.*, 54(2):208–217, Feb. 2010.
- [71] S. Warshall. A theorem on boolean matrices. 1962.

- [72] V. V. Williams and R. Williams. Subcubic equivalences between path, matrix and triangle problems. In *Proceedings of the 2010 IEEE 51st Annual Symposium on Foundations of Computer Science*, FOCS '10, pages 645–654, Washington, DC, USA, 2010. IEEE Computer Society.
- [73] N. Wirth. What can we do about the unnecessary diversity of notation for syntactic definitions? *Communications of the ACM*, 20(11):822–823, 1977.
- [74] D. Yan, G. Xu, and A. Rountev. Demand-driven context-sensitive alias analysis for java. In *Proceedings of the 2011 International Symposium on Software Testing and Analysis*, ISSTA '11, pages 155–165, New York, NY, USA, 2011. Association for Computing Machinery.
- [75] M. Yannakakis. Graph-theoretic methods in database theory. In *Proceedings of the ninth ACM SIGACT-SIGMOD-SIGART symposium on Principles of database systems*, pages 230–242. ACM, 1990.
- [76] H. Yu. An improved combinatorial algorithm for boolean matrix multiplication. In M. M. Halldórsson, K. Iwama, N. Kobayashi, and B. Speckmann, editors, *Automata, Languages, and Programming*, pages 1094–1105, Berlin, Heidelberg, 2015. Springer Berlin Heidelberg.
- [77] Q. Zhang, X. Xiao, C. Zhang, H. Yuan, and Z. Su. Efficient subcubic alias analysis for c. *SIGPLAN Not.*, 49(10):829–845, Oct. 2014.
- [78] X. Zhang, Z. Feng, X. Wang, G. Rao, and W. Wu. Context-free path queries on rdf graphs. In P. Groth, E. Simperl, A. Gray, M. Sabou, M. Krötzsch, F. Lecue, F. Flöck, and Y. Gil, editors, *The Semantic Web – ISWC 2016*, pages 632–648, Cham, 2016. Springer International Publishing.
- [79] X. Zheng and R. Rugina. Demand-driven alias analysis for c. In *Proceedings of the 35th Annual ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages*, POPL '08, pages 197–208, New York, NY, USA, 2008. ACM.
- [80] R. Zier-Vogel and M. Domaratzki. Rna pseudoknot prediction through stochastic conjunctive grammars. *Computability in Europe 2013. Informal Proceedings*, pages 80–89, 2013.
- [81] U. Zwick. A slightly improved sub-cubic algorithm for the all pairs shortest paths problem with real edge lengths. In *International Symposium on Algorithms and Computation*, pages 921–932. Springer, 2004.

- [82] В. Арлазаров, Е. Диниц, М. Кронрод, and И. Фараджев. Об экономном построении транзитивного замыкания ориентированного графа. *Докл. АН СССР*, 194(3):487–488, 1970.
- [83] Н. Вавилов. КОНКРЕТНАЯ ТЕОРИЯ ГРУПП, 2005. Дата доступа: 29 июня 2021 г.
- [84] Н. Вавилов. КОНКРЕТНАЯ ТЕОРИЯ КОЛЕЦ, 2006. Дата доступа: 29 июня 2021 г.