

МИНИСТЕРСТВО НАУКИ И ВЫСШЕГО ОБРАЗОВАНИЯ
РОССИЙСКОЙ ФЕДЕРАЦИИ
ВОЛЖСКИЙ ПОЛИТЕХНИЧЕСКИЙ ИНСТИТУТ (ФИЛИАЛ)
ФЕДЕРАЛЬНОГО ГОСУДАРСТВЕННОГО БЮДЖЕТНОГО ОБРАЗОВАТЕЛЬНОГО
УЧРЕЖДЕНИЯ ВЫСШЕГО ОБРАЗОВАНИЯ
«ВОЛГОГРАДСКИЙ ГОСУДАРСТВЕННЫЙ ТЕХНИЧЕСКИЙ УНИВЕРСИТЕТ»

Исследование операций: конспект лекций (часть 2)

Электронное учебное пособие



Волжский

2021

УДК 519.8(07)
УДК 22я73
И 889

Рецензенты:

доктор экономических наук, заведующий кафедрой теоретической
экономики и экономической безопасности ВИАПП, доцент

Орехова Е.А.,

доцент кафедры физики, математики и информатики ВолгГМУ,

канд.физ.-мат. наук

Мещерякова Н.Е.

Издается по решению редакционно-издательского совета
Волгоградского государственного технического университета

Алпатов, А.В.

Исследование операций : конспект лекций (часть2) [Электрон-
ный ресурс] : учебное пособие / А.В. Алпатов ; Министерство науки и
высшего образования Российской Федерации, ВПИ (филиал) ФГБОУ
ВО ВолгГТУ. – Электрон. текстовые дан. (1 файл: 1,16МБ). – Волж-
ский, 2021. – Режим доступа: <http://lib.volpi.ru>. – Загл. с титул. экрана.

ISBN 978-5-9948-4051-1

В учебно-методическом пособии приведены теоретические сведения об ана-
литических и численных методах решения задач оптимизации, а также примеры
решения указанных задач, которые могут быть использованы студентами при
подготовке к практическим занятиям и при выполнении контрольной работы.

Предназначено для студентов, обучающихся по направлению подготов-
ки 09.03.04 Программная инженерия в рамках курса «Исследование операций».

Ил. 23, табл. 6, библиограф.: 15 назв.

ISBN978-5-9948-4051-1

© Волгоградский государственный
технический университет, 2021

© Волжский политехнический
институт, 2021

СОДЕРЖАНИЕ

ВВЕДЕНИЕ.....	4
ГЛАВА 1. МНОГОКРИТЕРИАЛЬНАЯ ОПТИМИЗАЦИЯ.....	5
ГЛАВА 2. СЕТЕВОЕ ПЛАНИРОВАНИЕ И УПРАВЛЕНИЕ	18
ГЛАВА 3. ОПТИМИЗАЦИЯ НА ГРАФАХ	39
ГЛАВА 4. ДИНАМИЧЕСКОЕ ПРОГРАММИРОВАНИЕ.....	46
ГЛАВА 5. СИСТЕМЫ МАССОВОГО ОБСЛУЖИВАНИЯ	50
ГЛАВА 6. ОСНОВЫ ТЕОРИИ ИГР	76
ГЛАВА 7. ПРИНЯТИЕ РЕШЕНИЙ В УСЛОВИЯХ	93
НЕОПРЕДЕЛЕННОСТИ И РИСКА	93
Библиографический список.....	103

ВВЕДЕНИЕ

Учебное пособие «Исследование операций: конспект лекций (часть 2)» было разработано в соответствии с образовательным стандартом и предназначено для студентов, обучающихся по направлению подготовки 09.03.04 Программная инженерия.

Данное учебное пособие может использоваться студентами для подготовки к практическим и лабораторным занятиям, а также к промежуточной аттестации по дисциплине «Исследование операций». Изложение теории сопровождается примерами, разбором типовых задач. Учебное пособие разделено на семь глав: многокритериальная оптимизация, сетевое планирование и управление, оптимизация на графах, динамическое программирование, системы массового обслуживания, теория игр и принятие решений в условиях неопределенности и риска.

Исследование операций – дисциплина, занимающаяся разработкой и применением методов нахождения оптимальных решений на основе математического моделирования, статистического моделирования и различных эвристических подходов в различных областях человеческой деятельности.

Для успешного освоения дисциплины «Исследование операций» необходимо знать материал следующих дисциплин: математического анализа, линейной алгебры, теории вероятностей и математической статистики, дискретной математики, экономической теории.

ГЛАВА 1. МНОГОКРИТЕРИАЛЬНАЯ ОПТИМИЗАЦИЯ

1.1. Постановка задачи многокритериальной оптимизации

Известно, что экономическая эффективность производства количественно измеряется системой экономических показателей. Так, в промышленном производстве важными показателями эффективности являются: рост производительности труда, прибыль, рентабельность и др. Максимальное значение одного из показателей еще не означает, что предприятие работает лучше. Только система показателей может характеризовать эффективность всего производства.

Определение 1.1. Задачи, решаемые с учетом множества показателей или критериев, носят название *многокритериальных (многоцелевых)*.

Обозначим i -й частный критерий через $z_i(x)$, а область допустимых решений через Q . Учтем, что изменением знака функции всегда можно свести задачу минимизации к задаче максимизации, и наоборот, мы можем сформулировать кратко задачу многокритериальной оптимизации следующим образом:

$$Z(x) = \begin{pmatrix} z_1(x) \\ z_2(x) \\ \vdots \\ z_m(x) \end{pmatrix} \rightarrow \max, x \in Q \quad (1.1)$$

Еще один пример – выбор инвестиционного решения, когда хочется получить максимальный доход (или доходность) при наименьшем риске.

Разумеется, решения, которое одновременно удовлетворяло бы всем противоречивым требованиям, как правило, не существует. Но математика может помочь и при решении таких задач. Помощь эта состоит не в нахождении несуществующего решения, одновременно обращающего все критерии в максимум, а в отбрасывании заведомо плохих решений.

В идеальном случае в задаче (1.1) можно вести поиск такого решения, которое принадлежит пересечению множеств оптимальных решений всех однокритериальных задач. Но указанное пересечение обычно оказывается пустым множеством, и потому приходится рассматривать *переговорное множество* – множество допустимых решений, оптимальных по Парето.

1.2. Оптимизация по Парето

Изучение понятия оптимальности по Парето начнем с частного случая двух критериев. Обозначим буквой Y некоторую обобщенную характеристику произвольной инвестиционной операции, которую назовем *эффективностью* операции (в качестве Y можно взять доход, доходность в процентах от вложенной суммы, доходность в процентах годовых, внутреннюю норму доходности и т.п.). Часто невозможно заранее точно предсказать эффективность той или иной операции, и такие операции рассматривают как случайные величины. При этом в качестве *ожидаемой эффективности* такой инвестиционной операции используют математическое ожидание $M(Y)$ случайной величины Y .

Под *риском* инвестиционных операций мы понимаем отклонение реальных значений эффективности инвестиционной операции от прогнозируемой эффективности (как в меньшую сторону, так и в большую). Если Y – случайная эффективность инвестиционной операции, и в качестве ожидаемой эффективности операции мы выбрали математическое ожидание $M(Y)$ случайной величины Y , то в качестве *меры риска* операции естественно взять среднее квадратическое отклонение:

$$\sigma(Y) = \sqrt{D(Y)} \quad (1.2)$$

(здесь $D(Y)$ – дисперсия случайной величины Y).

Знание только математических ожиданий и средних квадратичных отклонений случайных величин довольно-таки важно при анализе группы случайных величин, оно помогает выбрать из множества случайных величин оптимальные по Парето, отбросив заведомо «плохие».

Пусть на финансовом рынке существует возможность осуществить несколько инвестиционных операций, ожидаемые эффективности и риски которых известны и равны соответственно: $M(Y_1), M(Y_2), \dots, M(Y_n)$; $\sigma(Y_1), \sigma(Y_2), \dots, \sigma(Y_n)$.

Определение 1.2. Говорят, что i -я операция *доминирует* j -ю, если:

$$\begin{cases} M(Y_i) \geq M(Y_j), \\ \sigma_i < \sigma_j \end{cases} \quad (1.3)$$

Или:

$$\begin{cases} M(Y_i) > M(Y_j), \\ \sigma_i \leq \sigma_j \end{cases} \quad (1.4)$$

Определение 1.3. Операция называется *оптимальной по Парето*, если не существует операций, которые бы ее доминировали.

Пример 1.1. Инвестор рассматривает четыре инвестиционные операции со случайными эффективностями, описываемыми случайными величинами Y_1, Y_2, Y_3, Y_4 с рядами распределения:

Y_1	2	4	6	8	Y_2	2	3	5	12
p	1/6	1/2	1/6	1/6	p	1/2	1/6	1/6	1/6
Y_3	3	4	8	10	Y_4	2	3	4	8
p	1/6	1/6	1/6	1/2	p	1/3	1/3	1/6	1/6

Необходимо определить, какие из этих операций оптимальны по Парето.

Решение:

Ожидаемые эффективности и риски равны соответственно:

$$M(Y_1) = 4,67, M(Y_2) = 4,33, M(Y_3) = 7,50, M(Y_4) = 3,67; \sigma(Y_1) = 0,58, \sigma(Y_2) = 3,59, \sigma(Y_3) = 2,93; \sigma(Y_4) = 2,05.$$

Нанесем точки () на единый график (рис. 1). i -я операция доминирует j -ю, если точка, соответствующая i -й операции, находится на графике правее и ниже точки, соответствующей j -й операции.

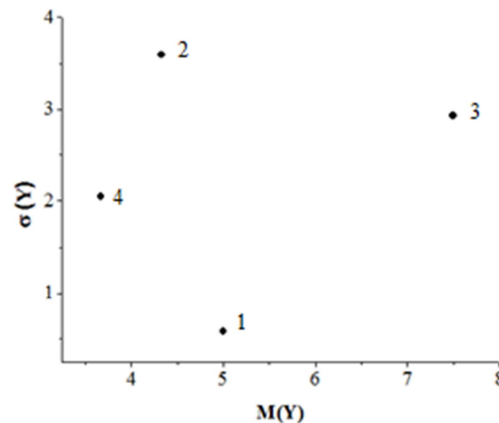


Рис. 1.1.

Видно, что первая операция доминирует вторую и четвертую, третья операция доминирует вторую. При этом первая операция не доминирует третью, а третья не доминирует первую. Первая и третья операции, таким образом, оптимальны по Парето.

Отметим, что операции, оптимальные по Парето, не обязательно являются «самыми лучшими» (и даже просто «хорошими») – эти операции не являются худшими. Выбор операций среди оптимальных по Парето осуществляется на основе склонности лица, принимающего соответствующее решение, к риску. В некоторых ситуациях предпочтительной оказывается операция, в которой ожидаемая эффективность вообще отрицательна. Например, если перед нами стоит выбор из двух операций:

1. потерять 1 руб.;

2. с вероятностью 0,5 получить 1000000 руб. и с вероятностью 0,5 потерять 100 000 руб.

Обе эти операции окажутся оптимальными по Парето ($M(Y_1) = -1$, $\sigma_1 = 0$, $M(Y_2) = 450000$, $\sigma_2 = 450\,000$), но, скорее всего, мы склонимся к выбору первой операции, несмотря на то, что ожидаемый доход по ней составляет отрицательное число (–1 руб.), тогда как ожидаемый доход от исполнения второй операции составляет 450000 руб. – слишком велик риск.

1.3. Оптимизация по Парето в случае нескольких критериев

Пусть необходимо выбрать несколько альтернатив на основе m критериев. Критерии выбора альтернатив задаются вектором оценки $\bar{f} = (f_1, f_2, \dots, f_m)$, где f_j – оценка возможного решения по критерию с номером j . В результате этого сравнение любых двух решений заменяется сравнением их векторных оценок. Сравнение векторных оценок осуществляется на основании принципа доминирования по Парето.

– Векторную оценку $\bar{f} = (f_1, f_2, \dots, f_m)$ называют доминирующей по Парето векторную оценку $\bar{g} = (g_1, g_2, \dots, g_m)$, если для всех $j=1, 2, \dots, m$, выполняется неравенство $f_j \geq g_j$, причем, хотя бы для одного значения j последнее неравенство строгое. Например, сравним две векторные оценки:

$$\bar{f} = \begin{pmatrix} 2 \\ 5 \\ 4 \\ 1 \end{pmatrix}, \quad \bar{g} = \begin{pmatrix} 3 \\ 2 \\ 7 \\ 4 \end{pmatrix}, \quad \bar{f} \geq \bar{g} = \begin{pmatrix} 0 \\ 1 \\ 0 \\ 0 \end{pmatrix}$$

Значение последнего логического выражения показывает, что в оценке $\bar{f}$ доминирование оценки $\bar{g}$ осуществляется только по второму критерию и, следовательно, векторная оценка $\bar{f}$ не доминирует векторную оценку $\bar{g}$. Для двух других векторных оценок:

$$\bar{f} = \begin{pmatrix} 4 \\ 5 \\ 4 \\ 3 \end{pmatrix}, \bar{g} = \begin{pmatrix} 3 \\ 2 \\ 2 \\ 2 \end{pmatrix}, \bar{f} \geq \bar{g} = \begin{pmatrix} 1 \\ 1 \\ 1 \\ 1 \end{pmatrix}$$

векторная оценка f доминирует векторную оценку g , так как значение логического выражения $f \geq g$ принимает значение «истина» для всех компонент сравниваемых оценок. Очевидно также, что если значение логического выражения $g \geq f$ принимает значение «ложь», то это значит, что реализуется противоположное доминирование, т.е. векторная оценка f доминирует векторную оценку g .

В результате сравнения решений все доминируемые (худшие по всем критериям) альтернативы отбрасываются, а все прочие (недоминируемые) оставляются. Они образуют множество недоминируемых альтернатив или множество Парето оптимальных решений. Векторная оценка из этого множества называется Парето–оптимальной, если в нём не существует никакой другой векторной оценки, доминирующей по Парето данную оценку.

Парето–оптимальность оценки означает, что она не может быть улучшена ни по одному из критериев, составляющих векторную оценку, без ухудшения по какому-либо другому критерию, входящему в векторную оценку.

Пример 1.2. Пусть имеется $N = 5$ возможных вариантов автомобильной дороги между двумя населёнными пунктами и требуется выбрать наиболее полезный вариант. Полезность вариантов будем оценивать по $M = 4$ критериям:

- $K1$ – длина трассы;
- $K2$ – потенциальная экономическая эффективность эксплуатации;
- $K3$ – стоимость технического обслуживания;
- $K4$ – продолжительность строительства.

В реальных условиях оценки каждого варианта трассы по каждому критерию осуществляются по результатам технико-экономического анализа. Мы примем гипотетические оценки в условных единицах, и результаты представим в таблице (матрице).

$$G1 = \begin{pmatrix} 4 & 8 & 4 & 2 & 5 & 8 \\ 3 & 8 & 2 & 8 & 7 & 7 \\ 6 & 7 & 6 & 4 & 4 & 3 \\ 4 & 9 & 5 & 8 & 2 & 6 \end{pmatrix}.$$

В каждом столбце в этой матрицы приведены оценки варианта трассы по заданным критериям. Предположим, что ищется вариант, у которого оценки по всем критериям максимальны. Сравнивая варианты по принципу доминирования по Парето, убеждаемся, что имеется единственный вариант (второй), который доминирует (превосходит) все остальные по каждому из критериев. Это значит, что Парето-оптимальное множество состоит из одного элемента, который и является оптимальным. Это второй вариант трассы.

Но такая ситуация на практике встречается редко и поэтому рассмотрим другой более типичный вариант, представленный матрицей:

$$G = \begin{pmatrix} 3 & 3 & 3 & 3 & 5 \\ 4 & 5 & 3 & 2 & 6 \\ 3 & 2 & 4 & 4 & 6 \\ 6 & 1 & 2 & 1 & 2 \end{pmatrix}$$

На основании принципа доминирования по Парето найдем множество Парето-оптимальных решений. Номера сравниваемых альтернатив будем записывать как верхний индекс.

$$G = \begin{pmatrix} 3 & 3 & 3 & 3 & 5 \\ 4 & 5 & 3 & 2 & 6 \\ 3 & 2 & 4 & 4 & 6 \\ 6 & 1 & 2 & 1 & 2 \end{pmatrix}, \quad G^{<1>} \geq G^{<2>} = \begin{pmatrix} 1 \\ 0 \\ 1 \\ 1 \end{pmatrix}, \quad G^{<1>} \geq G^{<3>} = \begin{pmatrix} 1 \\ 1 \\ 0 \\ 1 \end{pmatrix}, \quad G^{<1>} \geq G^{<4>} = \begin{pmatrix} 1 \\ 1 \\ 0 \\ 1 \end{pmatrix},$$

$$G^{<1>} \geq G^{<5>} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 1 \end{pmatrix}, \quad G^{<2>} \geq G^{<3>} = \begin{pmatrix} 1 \\ 1 \\ 0 \\ 0 \end{pmatrix}, \quad G^{<2>} \geq G^{<4>} = \begin{pmatrix} 1 \\ 1 \\ 0 \\ 1 \end{pmatrix}, \quad G^{<2>} \geq G^{<5>} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}$$

$$G^{<3>} \geq G^{<4>} = \begin{pmatrix} 1 \\ 1 \\ 1 \\ 1 \end{pmatrix}, \quad G^{<3>} \geq G^{<5>} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 1 \end{pmatrix}, \quad G^{<4>} \geq G^{<5>} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}$$

На основании сравнения видим, что альтернатива $G^{<3>}$ доминирует альтернативу $G^{<4>}$, альтернатива $G^{<5>}$ доминирует альтернативы $G^{<2>}$ и $G^{<4>}$. Это означает, что альтернативы $G^{<2>}$ и $G^{<4>}$ должны быть исключены из дальнейшего рассмотрения, так как они не могут входить во множество Парето-оптимальных решений. Матрица PG Парето-оптимальных решений имеет вид:

$$PG = \begin{pmatrix} 3 & 3 & 5 \\ 4 & 3 & 6 \\ 3 & 4 & 6 \\ 6 & 2 & 2 \end{pmatrix}$$

В неё включены только первая, третья и пятая альтернативы исходной матрицы G .

Таким образом, Парето-оптимальных оценок, как правило, получается множество, и выделить среди них единственную оптимальную оценку без дополнительной информации не представляется возможным, так как любые два Парето-оптимальных решения не сравнимы относительно доминирования по Парето. Это значит, что для любых двух Парето-оптимальных оценок всегда найдутся такие два критерия, по одному из которых предпочтительнее первый исход, а по другому – второй. В такой ситуации для принятия решения применяют два варианта.

В первом варианте выбор оптимального решения из множества Парето-оптимальных осуществляет ЛПР и, следовательно, решение имеет

субъективный характер. Во втором варианте на основе дополнительной информации выполняется сужение множества Парето-оптимальных решений с помощью некоторых формальных и эвристических процедур.

1.4. Метод обобщенного критерия

Метод обобщенного критерия предлагает перейти от m частных критериев к так называемой *свертке критериев*:

$$w = \sum_{i=1}^m \alpha_i z_i(x) \rightarrow \max, x \in Q \quad (1.5)$$

где веса α_i определяют важность критериев.

Пример 1.3. Дана задача многокритериальной оптимизации:

$$z_1 = -x_1 + 2x_2 \rightarrow \max$$

$$z_2 = 2x_1 + x_2 \rightarrow \max$$

$$z_3 = x_1 - 3x_2 \rightarrow \max$$

$$\begin{cases} x_1 + x_2 \leq 6, \\ 1 \leq x_1 \leq 3, \\ 1 \leq x_2 \leq 4, \\ x_1 \geq 0, \\ x_2 \geq 0 \end{cases}$$

Решить задачу методом обобщенного критерия, считая веса критериев равными $\alpha_1 = 0,5$, $\alpha_2 = \alpha_3 = 0,25$.

Решение:

Найдем множество допустимых решений, решим систему линейных неравенств. На рисунке 1.2 пятиугольник ABCDE является решением системы ограничений.

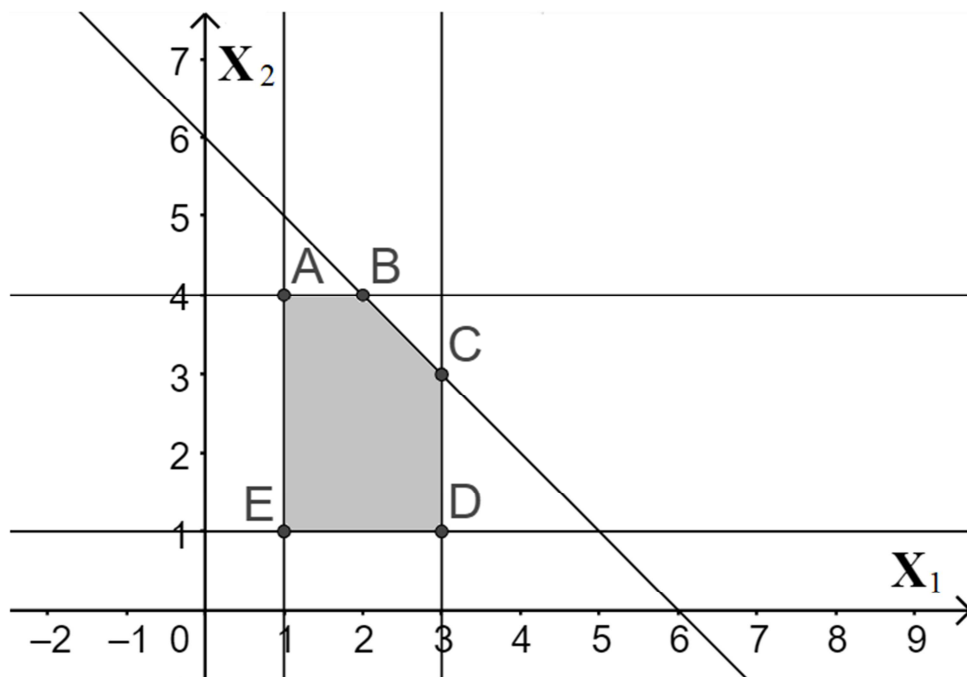


Рис. 1.2. Множество допустимых решений

Используя метод свертки критериев, из трех целевых функций получим одну:

Обобщённый критерий достигает максимум на множестве допустимых решений, как видно из рисунка 1.3, в точке B :

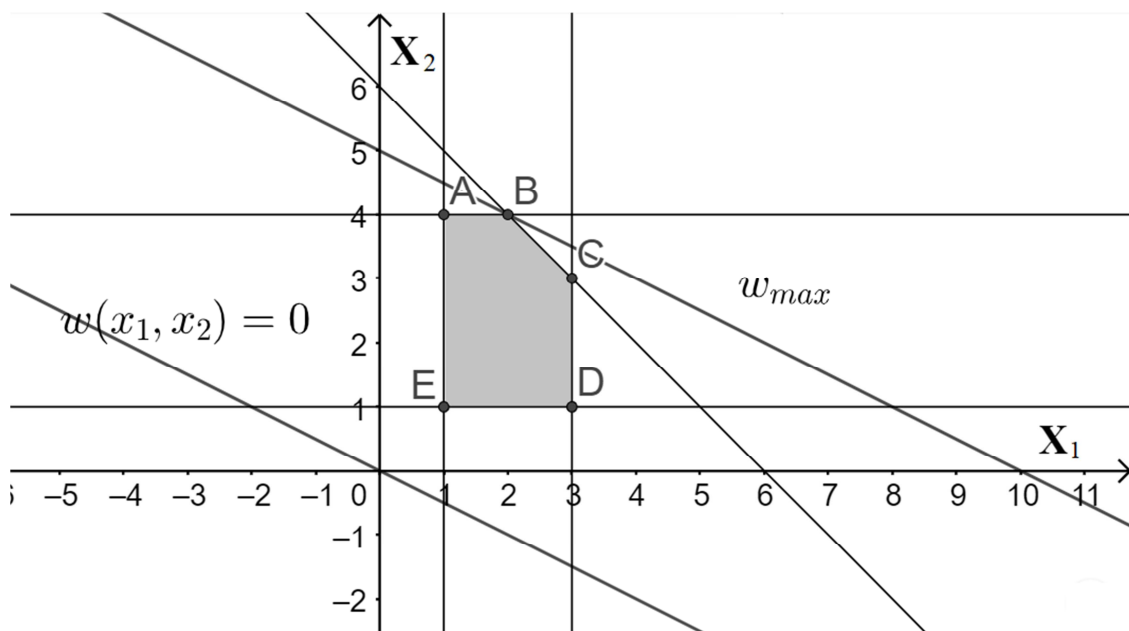


Рис. 1.3. Сверстка критериев

При этом:

1.4. Метод последовательных уступок

Метод последовательных уступок решения многокритериальных задач применяется в случае, когда частные критерии могут быть упорядочены в порядке убывающей важности.

Алгоритм метода последовательных уступок

1. Расположить критерии по их значимости (наиболее важным считается первым).
2. Решить задачу по первому критерию, , т.е. отыскать экстремальное значение целевой функции .
3. Сделать уступку по первому критерию, иными словами уменьшить величину до значения .

4. В задачу ввести дополнительное ограничение $z_1 \geq \alpha_1 z_1^*$.
5. Решить задачу по второму критерию $z_2 = z_2^*$.
6. Обратиться к пункту 3, сделать уступку для второго критерия $z_2 = \alpha_2 z_2^*$; $0 < \alpha_2 < 1$.
7. В задачу ввести дополнительное ограничение $z_2 \geq \alpha_2 z_2^*$.
8. Новую задачу, уже с двумя ограничениями, решить по третьему критерию и т.д.
9. Процесс итерации заканчивается, когда решение будет получено по всем критериям.

Пример 1.4. Решить задачу по двум критериям, считая первый наиболее предпочтительным. Его отклонение от максимального составляет не более 10 %.

$$z_1 = x_1 + 2x_2 \rightarrow \max, \quad (1.6)$$

$$z_2 = x_1 + x_2 \rightarrow \min, \quad (1.7)$$

$$\begin{cases} x_1 + 2x_2 \geq 6, \\ x_1 \leq 4, \\ x_2 \leq 5, \\ x_1 \geq 0, \\ x_2 \geq 0, \end{cases} \quad (1.8)$$

Используя первый критерий, находим область допустимых решений $ABCD$. Максимальное значение целевой функции достигается в точке $C(4,5)$, представленной на графике (рис. 13.4) $z_1^* = 4 + 2 \cdot 5 = 14$. Делаем уступку на 10% $z_1 = 0,9 \cdot 14 = 12,6$; получаем дополнительное ограничение $x_1 + 2x_2 > 12,6$ (на рис. 1.4 – прямая). С учетом дополнительного ограничения областью решения является треугольник ECK . Минимальное значение целевой функции z_2 достигается в точке $E(2,6; 5)$.

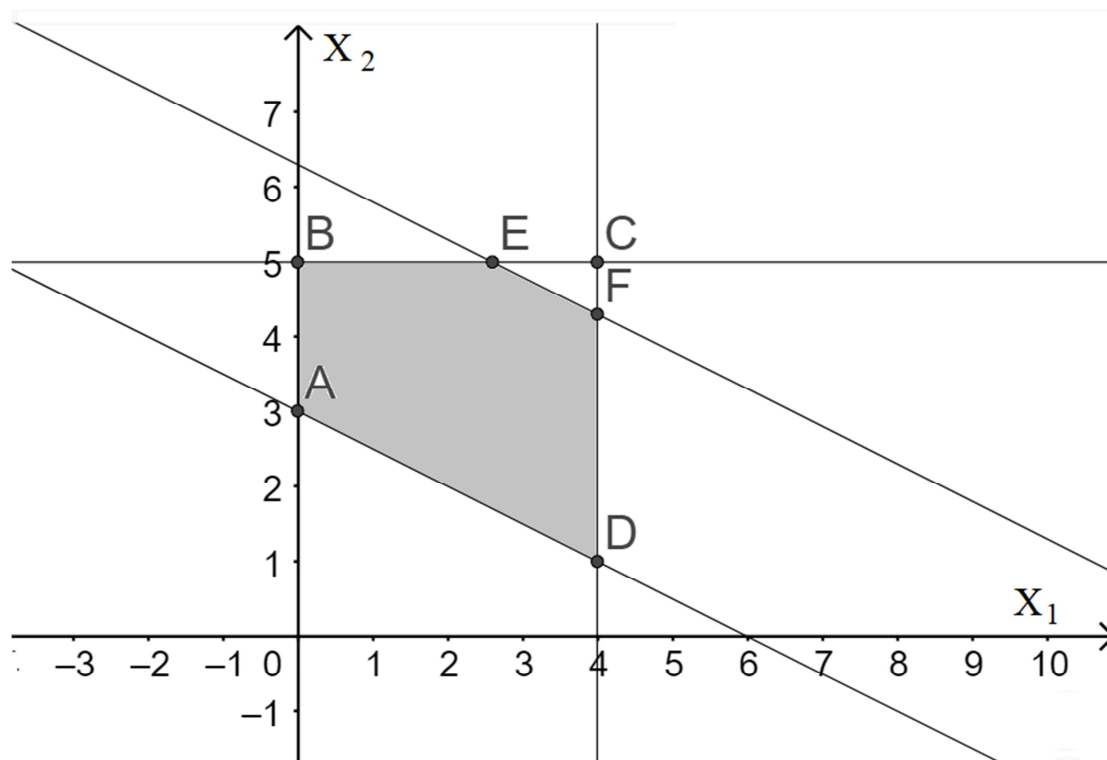


Рис. 1.4

Таким образом, план в точке $E(2,6; 5)$ будет наиболее эффективным при данных условиях:

ГЛАВА 2. СЕТЕВОЕ ПЛАНИРОВАНИЕ И УПРАВЛЕНИЕ

2.1. Понятие сетевой модели

Сетевое планирование и управление (СПУ) основано на моделировании с помощью сетевого графика и представляет собой совокупность расчетных методов, организационных и контрольных мероприятий по планированию и управлению комплексом работ.

Система СПУ позволяет:

1. формировать календарный план реализации некоторого комплекса работ;
2. выявлять и мобилизовать резервы времени, трудовые материальные и денежные ресурсы;
3. осуществлять управление комплексом работ по принципу «ведущего звена» с прогнозированием и предупреждением возможных срывов в ходе работы;
4. повышать эффективность управления в целом при четком распределении ответственности между руководителями разных уровней и исполнителями работ.

Диапазон применения СПУ довольно широк: от задач, касающихся деятельности отдельных лиц, до проектов, в которых участвуют сотни организаций и десятки тысяч людей (например, разработка и создание крупного территориально-промышленного комплекса).

Определение 2.1. Под *комплексом работ* (комплексом операций или проектом) будем понимать всякую задачу, для выполнения которой необходимо осуществить достаточно большое количество разнообразных работ.

Определение 2.2. *Сетевая модель* представляет собой план выполнения некоторого комплекса взаимосвязанных работ (операций), заданного в

специфической форме сети, графическое изображение которой называется сетевым графиком.

Сетевые графики представляют собой оргграфы без контуров, дугам или вершинам которых приписаны некоторые величины.

Отличительной особенностью сетевой модели является четкое определение всех временных взаимосвязей предстоящих работ.

Главными элементами сетевой модели являются события и работы.

Определение 2.3. Под термином *работа* понимают:

- во-первых, это действительная работа – протяженный по времени процесс, требующий затрат ресурсов. Каждая действительная работа должна быть конкретной, четко описанной и иметь ответственного исполнителя;
- во-вторых, это ожидание – протяженный по времени процесс, не требующий затрат труда;
- в-третьих, это зависимость, или фиктивная работа – логическая связь, не требующая затрат труда, материальных ресурсов или времени. Она указывает, что возможность одной работы непосредственно зависит от результатов другой. Продолжительность фиктивной работы принимается равной нулю.

Определение 2.4. *Событие* – это момент завершения какого-либо процесса, отражающий отдельный этап выполнения проекта.

Событие может являться частным результатом отдельной работы или суммарным результатом нескольких работ. Событие может совершиться только тогда, когда закончатся все работы, ему предшествующие. Последующие работы могут начаться только тогда, когда событие совершиться.

Отсюда двойственный характер события:

- для всех непосредственно предшествующих ему работ оно является конечным;
- для всех непосредственно следующих за ним – начальным.

При этом предполагается, что событие не имеет продолжительности и свершается как бы мгновенно.

Среди событий сетевой модели выделяют исходное и завершающее события.

Определение 2.5. *Исходное событие* не имеет предшествующих работ и событий, *завершающее событие* не имеет последующих работ и событий.

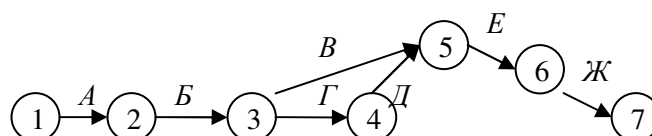


Рис. 2.1

Примерная задача сетевого планирования (см. рис. 2.1):

A – сформировать проблему исследования;

B – построить математическую модель изучаемого объекта;

B – собрать информацию;

Г – выбрать метод решения задачи;

Д – построить и отладить программу для ЭВМ;

E – рассчитать оптимальный план;

Ж – передать результаты расчета заказчику.

Цифрами обозначены номера событий, к которым приводит выполнение соответствующих работ.

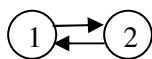
Анализ:

- работы *B* и *Г* можно начать выполнять независимо одна от другой только после совершения события 3, т.е. когда выполнены работы *A*, *B*;
- работу *Д* после совершения события 4, когда выполнены работы *A*, *B*, и *Г*;
- работу *E* можно выполнить после наступления события 5, т.е. при выполнении всех предшествующих работ *A*, *B*, *B*, *Г* и *Д*.

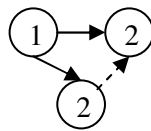
Если в сетевой модели не представлены числовые оценки, то такая сеть называется структурной.

При построении сетевого графика необходимо соблюдать ряд правил:

1. Вначале планируемый процесс разбивается на отдельные работы, составляется перечень работ и событий, продумываются их логические связи и последовательность выполнения, работы закрепляются за ответственными исполнителями. Затем составляется сетевой график.
2. В сетевой модели не должно быть «тупиковых событий», т.е. событий, из которых не выходит ни одна работа, за исключением завершающего события.
3. В сетевом графике не должно быть «хвостовых событий» (кроме исходного), которым не предшествует хотя бы одна работа.
4. В сети не должно быть замкнутых контуров и петель, т.е. путей, соединяющих некоторые события с ними самими.
5. Любые два события должны быть непосредственно связаны не более чем одной работой стрелкой.



– если такие работы встречаются, то рекомендуется ввести «фиктивное событие 2'» и «фиктивную работу», при этом одна из



параллельных работ замыкается на – это фиктивное событие 2'.

Фиктивная работа на графике изображается пунктирными линиями.

6. В сети рекомендуется иметь одно исходное и одно завершающее событие.

Если в сетевом графике это не так, то добиться желаемого можно путем введения фиктивных событий, работ. Фиктивные работы можно вводить в следующих случаях.

Один из них – отражение зависимости событий, не связанных с реальными работами.

Другой случай – неполная зависимость работ, кроме того, фиктивные работы могут вводиться для отражения реальных отсрочек и ожидания. В отличие от предыдущих случаев, здесь фиктивная работа имеет протяженность во времени.

Определение 2.6. *Упорядочение сетевого графика* – заключается в таком расположении событий и работ, при котором для любой работы предшествующее ей событие расположено левее и имеет меньший номер по сравнению с завершающим эту работу событием.

Другими словами, в упорядоченном сетевом графике все работы-стрелки направлены слева направо: от событий с меньшими номерами к событиям с большими номерами.

Одно из важнейших понятий сетевого графика – понятие пути.

Определение 2.7. *Путь* – любая последовательность работ, в которой конечное событие каждой работы совпадает с начальным событием следующей за ней работы.

Определение 2.8. *Полный путь L* – любой путь, начало которого совпадает с исходным событием сети, а конец – с завершающим.

Определение 2.9. Наиболее продолжительный полный путь в сетевом графике называется **критическим**.

Критическими также называются работы и события, расположенные на этом пути. Работы этого пути определяют общий цикл завершения всего комплекса работ, планируемых при помощи сетевого графика. И для сокращения продолжительности проекта необходимо в первую очередь сокращать продолжительность работ, лежащих на критическом пути.

Классический вид сетевого графика – это сеть, вычерченная без масштаба времени. Поэтому сетевой график, хотя и дает четкое представление о порядке следования работ, но недостаточно нагляден для определения тех работ, которые должны выполняться в каждый момент времени. В свя-

зи с этим небольшой проект после упорядочения сетевого графика рекомендуется дополнить линейной диаграммой проекта.

При построении линейной диаграммы каждая работа изображается параллельным оси времени отрезком, длина которого равна продолжительности этой работы. При наличии фиктивной работы нулевой продолжительности (в рассматриваемой сети ее нет) она изображается точкой. События i и j , начало и конец работы (i, j) помещают соответственно в начале и конце отрезка. Отрезки располагают один за другим, снизу вверх в порядке возрастания индекса i , а при одном и том же i – в порядке возрастания индекса j .

По линейной диаграмме проекта можно определить критическое время, критический путь, а также резервы времени всех работ. Так, критическое время комплекса работ равно координате на оси времени самого правого конца всех отрезков диаграммы.

Таблица 2.1

Временные параметры сетевых графиков

Элемент сети, характеризующийся параметром	Наименование параметра	Условное обозначение параметра
Событие i	Ранний срок свершения события	$t_p(i)$
	Поздний срок свершения события	$t_n(i)$
	Резерв времени события	$R(i)$
Работа (i, j)	Продолжительность работы	$t(i, j)$
	Ранний срок начала работы	$t_{pn}(i, j)$
	Ранний срок окончания работы	$t_{po}(i, j)$
	Поздний срок начала работы	$t_{nn}(i, j)$
	Поздний срок окончания работы	$t_{no}(i, j)$
	Полный резерв времени работы	$R_n(i, j)$
	Частный резерв времени первого вида	$R_I(i, j)$

	Частный резерв времени второго вида или свободный резерв времени	$R_c(i,j)$
	Независимый резерв времени	$R_n(i,j)$
Путь L	Продолжительность пути	$t(L)$
	Продолжительность критического пути	$t_{кр}$
	Резерв времени пути	$R(L)$

1. Ранний (или ожидаемый) срок совершения i -го события $t_p(i)$ определяется продолжительностью максимального пути предшествующему этому событию:

$$t_p(i) = \max_{L_{ni}} t(L_{ni}) \quad (2.1)$$

Для расчета раннего срока j -го события используют формулу:

$$t_p(j) = t_p(i) + t(i, j) \quad (2.2)$$

Если событие j имеет несколько предшествующих событий, то ранний срок события определяют по формуле:

$$t_p(j) = \max_i \{t_p(i) + t(i, j)\} \quad (2.3)$$

2. Поздний допустимый срок $t_n(i)$ наступления события – это максимальный срок, который не нарушает поздних допустимых сроков наступления следующих за ним событий.

Поздний допустимый срок события i равен разности между критическим временем проекта $t_{кр}$ и максимальной длиной пути из i -го события в n – это время необходимое для выполнения всех работ, происходящих после наступления события;

– если из события выходит одна стрелка;

$$t_n(i) = t_n(j) - t(i, j) \quad (2.4)$$

– если из события выходит несколько стрелок

$$t_n(i) = \min_j \{t_n(j) - t(i, j)\} \quad (2.5)$$

Резерв времени $R(i)$ i -го события определяется как разность между поздним и ранним сроками его свершения:

$$R(i) = t_n(i) - t_p(i) \quad (2.6)$$

Резерв времени показывает, на какой допустимый период времени можно задержать наступление этого события, не вызывая при этом увеличения срока выполнения комплекса работ.

Расчет ранних и поздних сроков событий дает еще один способ нахождения критического пути, у работ, лежащих на критическом пути, совпадают ранние и поздние сроки событий $t_p(i) = t_n(i)$.

Совпадение ранних и поздних сроков событий для определения критического пути является условием необходимым, но не достаточным (это условие не дает однозначного определения критического пути).

Любая задержка работ, лежащих не на критическом пути, нарушает сроки выполнения всего комплекса работ. Работы, не лежащие на критическом пути, допускают определенную задержку во времени выполнения, которая не сказывается на сроках выполнения всего комплекса работ.

Если сетевой график имеет единственный критический путь, то этот путь проходит через все критические события, т.е. события с нулевыми резервами времени. Если критических путей несколько, то выявление их с помощью критических событий может быть затруднено, так как через часть критических событий могут проходить как критические, так и не критические пути. В этом случае для определения критических путей рекомендуется использовать критические работы.

Отдельная работа может начаться (и окончиться) в ранние, поздние или другие промежуточные сроки. В дальнейшем при оптимизации графика возможно любое размещение работы в заданном интервале.

Ранний срок $t_{рн}(i, j)$ начала работы (i, j) совпадает с ранним сроком наступления начального (предшествующего) события i , т.е.:

$$t_{рн}(i, j) = t_p(i) \quad (2.7)$$

Тогда ранний срок $t_{ро}(i, j)$ окончания работы (i, j) определяется по формуле:

$$t_{ро}(i, j) = t_p(i) + t(i, j) \quad (2.8)$$

Ни одна работа не может закончиться позже допустимого позднего срока своего конечного события j . Поэтому поздний срок $t_{по}(i, j)$ окончания работы (i, j) определяется соотношением:

$$t_{по}(i, j) = t_n(j) \quad (2.9)$$

а поздний срок $t_{пн}(i, j)$ начала данной работы (i, j) – соотношением:

$$t_{пн}(i, j) = t_n(j) - t(i, j) \quad (2.10)$$

Резерв времени пути $R(L)$ определяется как разность между длиной критического и рассматриваемого пути:

$$R(L) = t_{кр} - t(L) \quad (2.11)$$

Он показывает, насколько в сумме могут быть увеличены продолжительности всех работ, принадлежащих этому пути. Если затянуть выполнение работ, лежащих на этом пути, на время большее чем $R(L)$, то критический путь переместится на путь L .

Можно определить четыре вида резерва времени выполнения работы (i, j) .

Полный резерв времени $R_n(i, j)$ работы (i, j) – показывает, на сколько можно увеличить время выполнения данной работы при условии, что срок выполнения комплекса работ не изменится. Полный резерв времени определяется по формуле:

$$R_n = t_n(j) - t_p(i) - t(i, j) \quad (2.12)$$

Полный резерв времени работы равен резерву максимального из путей, проходящего через данную работу. Этим резервом времени можно

располагать при выполнении данной работы, если ее начальное событие свершится в ранний срок, и можно допустить свершение конечного события в его самый поздний срок.

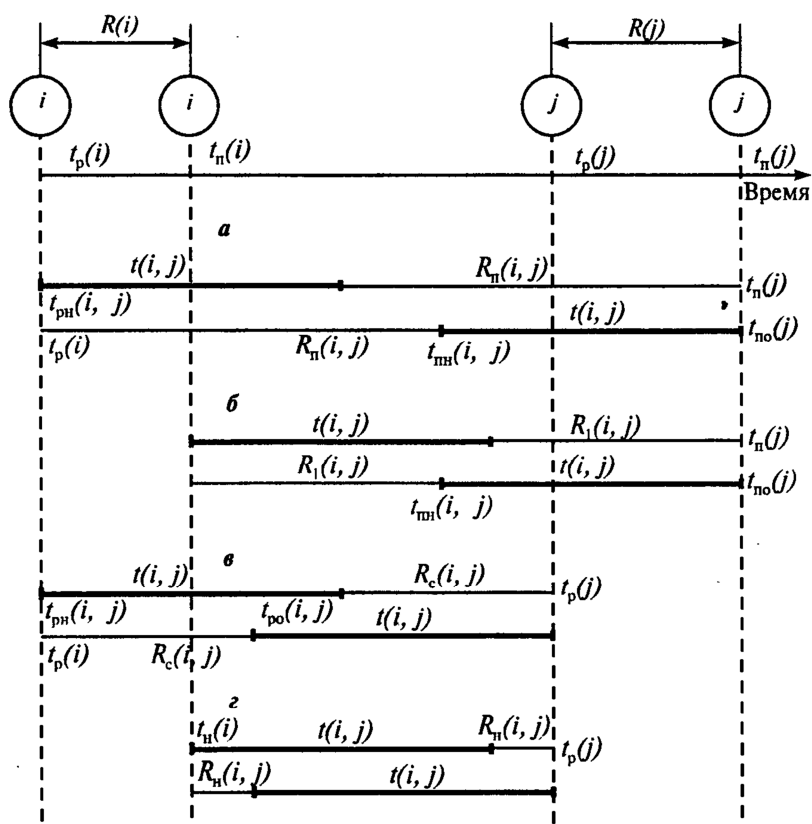


Рис. 2.2

Важной особенностью полного резерва времени работы является то, что он принадлежит не только этой работе, но и всем полным путям, проходящим через нее. При использовании полного резерва времени только для одной работы резервы времени остальных работ, лежащих на других путях, проходящих через эту работу, сократятся соответственно на величину использованного резерва. Остальные резервы времени работы являются частями ее полного резерва.

Частный резерв времени первого – гарантийный резерв, это часть полного резерва времени, на которую можно увеличить продолжительность работы, не изменив при этом позднего срока ее начального события.

Этим резервом можно располагать при выполнении данной работы в предположении, что ее начальное и конечное события свершаться в свой самый поздние сроки (см. рис. 2.2):

$$R_1(i, j) = t_{\pi}(j) - t_{\pi}(i) - t(i, j) \quad (2.13)$$

$$R_1(i, j) = R_{\pi}(i, j) - R(i) \quad (2.14)$$

Частный резерв времени второго вида $R_c(i, j)$ – свободный резерв времени, представляет собой часть полного резерва времени, на которую можно увеличить продолжительность работы, не изменив при этом раннего срока ее конечного события.

Этим резервом можно располагать при выполнении данной работы в предположении, что ее начальное и конечное события свершаться в свои самые ранние сроки:

$$R_c(i, j) = t_p(j) - t_p(i) - t(i, j) \quad (2.15)$$

$$R_c(i, j) = R_{\pi}(i, j) - R(j) \quad (2.16)$$

Свободным резервом времени можно пользоваться для предотвращения случайностей, которые могут возникнуть в ходе выполнения работ по ранним срокам их начала и окончания, то всегда будет возможность при необходимости перейти на поздние сроки начала и окончания работ.

Независимый резерв времени $R_n(i, j)$ – часть полного резерва времени, получаемая для случая, когда все предшествующие работы заканчиваются в поздние сроки, а все последующие работы начинаются в ранние сроки:

$$R_n(i, j) = t_p(j) - t_{\pi}(i) - t(i, j) \quad (2.17)$$

$$R_n(i, j) = R_{\pi}(i, j) - R(i) \quad (2.18)$$

Использование независимого резерва времени не влияет на величину резервов времени других работ. Независимые резервы стремятся использовать тогда, когда окончание предыдущей работы произошло в поздний допустимый срок, а последующие работы хотят выполнять в ранние сроки. Если величина независимого резерва равна нулю или положительна, то такая возможность есть. Если же величина $R_n(i, j)$ отрицательна, то этой

возможности нет, так как предыдущая работа еще не оканчивается, а последующая уже должна начаться. Поэтому отрицательное значение $R_n(i, j)$ не имеет реального смысла. А фактически независимый резерв имеют лишь те работы, которые не лежат на максимальных путях, проходящих через их начальные и конечные события.

Следует отметить, что резервы времени работы (i, j) , показанные на рисунке 2.2, могут состоять из двух временных отрезков, если интервал продолжительности работы $t(i, j)$ занимает промежуточную позицию между двумя его крайними положениями, изображенными на графиках.

Таким образом, если частный резерв времени первого вида может быть использован на увеличение продолжительности данной и последующих работ без затрат резерва времени предшествующих работ, а свободный резерв времени – на увеличение продолжительности данной и предшествующих работ без нарушения резерва времени последующих работ, то независимый резерв времени может быть использован для увеличения продолжительности только данной работы.

Работы, лежащие на критическом пути, так же, как и критические события, резервов времени не имеют.

Если на критическом пути лежит начальное событие i , то:

$$R_n(i, j) = R_1(i, j) \quad (2.19)$$

Если на критическом пути лежит конечное событие j , то:

$$R_n(i, j) = R_c(i, j) \quad (2.20)$$

Если на критическом пути лежат начальное и конечное события i и j , но сама работа не принадлежит этому пути, то:

$$R_n(i, j) = R_1(i, j) = R_c(i, j) = R_n(i, j) \quad (2.21)$$

Соотношения (7.19)-(7.21) можно использовать при проверке правильности расчетов резервов времени отдельных работ.

Замечание: если какой-либо резерв времени оказывается отрицательным, то его приравнивают к нулю.

2.2. Сетевое планирование в условиях неопределенности

При определении временных параметров сетевого графика мы исходим из предположения, что время выполнения каждой работы точно известно.

Однако на практике система СПУ обычно применяется для планирования сложных разработок, не имевших в прошлом никаких аналогов.

Чаще всего продолжительность работ по сетевому графику заранее не известна и может принимать лишь одно из ряда возможных значений. Другими словами, продолжительность работы (i, j) является случайной величиной, характеризующейся своим законом распределения, а значит, и своими числовыми характеристиками – математическим ожиданием, средним значением $\overline{t_{ij}}$ и дисперсией $D = \sigma^2(i, j)$.

Практически во всех системах СПУ принимается допущение, что распределение продолжительности работ обладает тремя свойствами:

- непрерывностью;
- унимодельностью, т.е. наличие единственного максимума у кривой распределения;
- двумя точками пересечения кривой распределения с осью Ox , имеющими неотрицательные абсциссы.

Кроме того, принимается, что распределение продолжительности работ обладает положительной асимметрией, т.е. максимум кривой смещен влево относительно медианы (линии, делящей площадь под кривой на две равные части). Распределение, как правило, более круто поднимается при удалении от минимального значения t и полного опускается при приближении к максимальному значению t .

Простейшим распределением с подобными свойствами является известное в математической статистике β -распределение. Анализ большого количества статистических данных (хронометража времени реализации

отдельных работ, нормативные данные и т.п.) показывает, что β -распределение можно использовать в качестве распределения продолжительности работ.

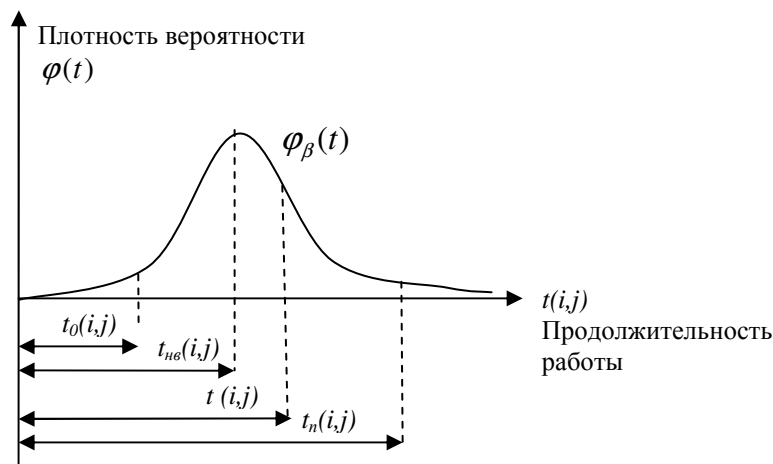


Рис. 2.3

Для определения числовых характеристик $\bar{t}_{ij}$ и $\sigma^2(i, j)$ данного распределения для работы (i, j) на основании опроса ответственных исполнителей проекта и экспертов определяют три временные оценки:

- $t_0(i, j)$ – оптимистическую оценку, т.е. продолжительности работы (i, j) при самых благоприятных условиях;
- $t_n(i, j)$ – пессимистическую оценку, т.е. продолжительности работы (i, j) при самых неблагоприятных условиях;
- $t_{н\phi}(i, j)$ – наиболее вероятную оценку продолжительности работы (i, j) при нормальных условиях.

Предположение о β -распределении продолжительности работы (i, j) позволяет получить оценки ее числовых величин:

$$\bar{t}(i, j) = \frac{t_0(i, j) + 4t_{н\phi}(i, j) + t_n(i, j)}{6}$$

$$\sigma^2(i, j) = \left[\frac{t_n(i, j) - t_0(i, j)}{6} \right]^2$$

Обычно оценить наиболее вероятное время выполнения работы $t_{нв}(i, j)$ довольно сложно, поэтому реально на практике используют упрощенную (менее точную) оценку средней продолжительности работы $\bar{t}_{ij}$ на основании лишь двух задаваемых временных оценок $t_0(i, j)$ и $t_n(i, j)$:

$$\bar{t}(i, j) = \frac{2t_0(i, j) + 3t_n(i, j)}{5}$$

Зная $\bar{t}_{ij}$ и $\sigma^2(i, j)$, можно определить временные параметры сетевого графика и оценить их надежность.

При достаточно большом количестве работ, принадлежащих пути L , и выполнении некоторых весьма общих условий можно применить центральную теорему Ляпунова, на основании которой можно утверждать, что общая продолжительность пути L имеет нормальный закон распределения со средним значением $\bar{t}(L)$, равное сумме средних значений продолжительности составляющих его работ $\bar{t}(i, j)$ и дисперсией $\sigma^2(L)$, равный сумме соответствующих дисперсий $\sigma^2(i, j)$:

$$\begin{aligned}\bar{t}(L) &= \sum_{i,j} \bar{t}(i, j) \\ \sigma^2(L) &= \sum_{i,j} \sigma^2(i, j)\end{aligned}$$

Сетевой график, в этом случае, представляет сеть не с детерминированными (фиксированными), а со случайными продолжительностями работ, и цифры над работами-стрелками указывают среднее значение $\bar{t}(i, j)$ продолжительности соответствующих операций, вычисленное по одной из приведенных выше формул, и известны все дисперсии $\sigma^2(i, j)$.

Следует отметить, что и в этом случае временные параметры сетевого графика:

- длина критического пути $t_{кр}$;
- ранние и поздние сроки свершения событий $\{t_p(i), t_p(j); t_n(i), t_n(j)\}$;

- резервы времени событий и работ определяются так же, как рассматривалось выше.

Но при этом необходимо учитывать, что эти параметры теперь будут являться средними значениями соответствующих случайных величин:

- средней длиной критического пути $\bar{t}_{кр}$;
- средним значением данного срока наступления события $\bar{t}_p(i)$;
- средним значением полного резерва времени работы $\bar{R}_n(i, j)$ и т.д.

Так, $\bar{t}_{кр} = 61$ суток будет означать, что длина критического пути лишь в среднем составляет – 61 сутки, а в каждом конкретном случае возможны заметные отклонения от среднего значения (причем, чем больше суммарная дисперсия продолжительности работ критического пути, тем более вероятны значительные по абсолютной величине отклонения).

Поэтому предварительный анализ сетей со случайными продолжительностями работ, как правило, не ограничивается расчетами временных параметров сети.

В этом случае необходимо оценить вероятность того, что срок выполнения проекта $t_{кр}$ не превзойдет заданного директивного срока T .

Полагая $t_{кр}$ случайной величиной, имеющий нормальный закон распределения, получим:

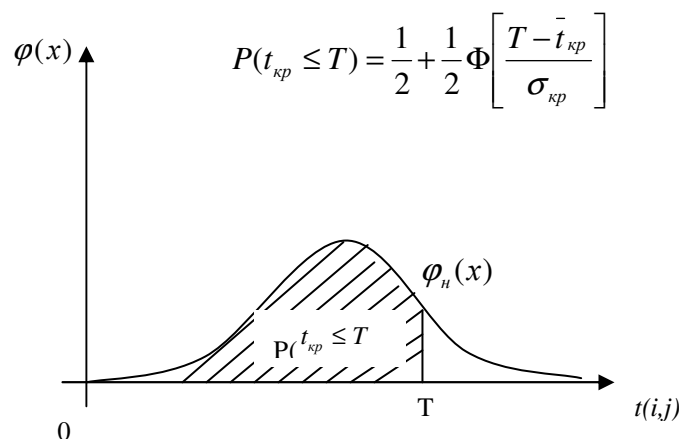


Рис. 2.4

Численно эта вероятность равна площади заштрихованной фигуры, где $\Phi(z) = \Phi\left(\frac{T - \bar{t}_{кр}}{\sigma_{кр}}\right)$ – значение интеграла вероятностей Лапласа, $\sigma_{кр}$ – среднее квадратическое отклонение длины критического пути $\sigma_{кр} = \sqrt{\sigma_{кр}^2}$.

$t_{кр}$ и $\sigma_{кр}^2$ определяются по формулам:

$$\bar{t}(L) = \sum_{i,j} t(i, j)$$

$$\sigma^2(L) = \sum_{i,j} \sigma^2(i, j)$$

Если $P(t_{кр} \leq T) \leq 0,3$ (мало), опасность срыва заданного срока выполнения комплекса работ велика, необходимо принятие дополнительных мер (перераспределение ресурсов по сети, пересмотр состава работ, событий и т.п.).

Если $P(t_{кр} \leq T) \approx 0,8$, то очевидно с достаточной степенью надежности можно прогнозировать выполнение проекта в установленный срок.

В некоторых случаях возможно решение обратной задачи: определение максимального срока выполнения проекта с заданной надежностью (вероятностью) β . В этом случае:

$$T = \bar{t}_{кр} + z_{\beta} \sigma_{кр}^2,$$

где z_{β} – нормированное отклонение случайной величины, определяемое с помощью функции Лапласа $\Phi(z_{\phi}) = \beta$.

Замечание:

1. На основании теоремы Ляпунова вывод о нормированном законе распределения случайной величины $\bar{t}_{кр}$ правомерен лишь для достаточно большого числа критических работ (отсюда низкая точность $P(t_{кр} \leq T)$).
2. Приведенный метод расчета имеет принципиальный недостаток оценки параметров сложных сетей с большим количеством работ. Это можно объяснить тем, что на практике нередки случаи, когда дисперсии

$\sigma^2(t)$ длин не критических (но близкие к критическому) существенно больше, чем $\sigma_{кр}^2$. Поэтому при изменении ряда условий в данном конкретном комплексе работ возможен переход к новым критическим путям, которые в расчете не учитываются.

3. Различия между событием с детерминированными случайными продолжительностями работ не следует путать с различием детерминированных и стохастических сетей. Последнее различие связано со структурой самой сети.

Рассмотренные выше сети являлись детерминированными. Работы в них могли характеризоваться не только детерминированными, но и случайными продолжительностями.

Однако встречаются проекты, когда на некоторых этапах тот или иной комплекс последующих работ зависит от неизвестного заранее результата. Какой из этих комплексов работ будет фактически выполняться, заранее неизвестно, и может быть предсказано с некоторой вероятностью.

Например: может быть предусмотрено несколько вариантов продолжения исследования в зависимости от полученных опытных данных или несколько вариантов строительства предприятий различной мощности по обработке сырья в зависимости от результатов разведки запасов этого сырья. Такие сети называются *стохастическими*.

В свою очередь, стохастические сети так же, как и детерминированные, могут характеризоваться детерминированными либо случайными продолжительностями.

2.3. Коэффициент напряженности работы

После нахождения критического пути и резерва времени работ и оценки вероятности выполнения проекта в заданный срок должен быть

проведен всесторонний анализ сетевого графика и приняты меры его оптимизации.

Анализ сетевого графика начинается с анализа топологии сети, включающего контроль построения сетевого графика, установление целесообразности выбора работ, степени их расчленения.

Затем проводятся классификация и группировка работ по величинам резервов. Необходимо учитывать, что величина полного резерва времени далеко не всегда может достаточно точно характеризовать, насколько напряженным является выполнение той или иной работы критического пути.

Определить степень трудности выполнения в срок каждой группы работ некритического пути можно с помощью коэффициента напряженности работ.

Коэффициентом напряженности K_n работы (i, j) называется отношение продолжительности несовпадающих (заключенных между одними и теми же событиями) отрезков пути, одним из которых является путь максимальной продолжительности, проходящий через данную работу, а другим – критический путь.

$$K_n(i, j) = \frac{t(L_{max}) - t'_{кр}}{t_{кр} - t'_{кр}} \quad (2.22)$$

где $t(L_{max})$ – продолжительность максимального пути, проходящего через работу (i, j) ;

$t_{кр}$ – продолжительность (длина) критического пути;

$t'_{кр}$ – продолжительность отрезка рассматриваемого пути, совпадающего с критическим путем.

Формулу для K_n можно привести к виду:

$$K_n(i, j) = 1 - \frac{R_n(i, j)}{t_{кр} - t'_{кр}} \quad (2.23)$$

здесь $R_n(i, j)$ – полный резерв времени работы (i, j) .

Коэффициент напряженности $K_n(i,j)$ может изменяться в пределах от 0 (для работ, у которых отрезки максимального из путей не совпадающие с критическим путем, состоит из фиктивных работ нулевой продолжительности) до 1 (для работ критического пути).

3. Оптимизация сетевого графика

Оптимизация сетевого графика представляет процесс улучшения организации выполнения комплекса работ с учетом срока его выполнения. Оптимизация проводится с целью сокращения критического пути, выравнивании коэффициентов напряженности работ, рационального использования ресурсов.

В первую очередь принимаются меры по сокращению продолжительности работ, находящихся на критическом пути. Это достигается:

- перераспределением всех видов ресурсов, как временных (использование резервов времени некритических путей), так и трудовых, материальных, энергетических (например, перевод части исполнителей, оборудования с некритических путей на работы критического пути); при этом перераспределение ресурсов должно идти, как правило, из зон, менее напряженных, в зоны, объединяющие наиболее напряженные работы;
- сокращением трудоемкости критических работ за счет передачи части работ на другие пути, имеющие резервы времени;
- параллельным выполнением работ критического пути;
- пересмотром топологии сети, изменением состава работ и структуры сети.

В процессе сокращения продолжительности работы критический путь может измениться, и в дальнейшем процесс оптимизации будет направлен на сокращение продолжительности работ нового критического пути и так

будет продолжаться до получения удовлетворительного результата. В идеале длина любого из полных путей может стать равной длине критического пути или, по крайней мере, пути критической зоны. Тогда все работы будут вестись с равным напряжением, а срок завершения проекта существенно сократиться.

Весьма эффективным является использование метода *статистического моделирования*, основанного на многократных последовательных изменениях продолжительности работ (в заданных пределах) и «проигрывание» на компьютере различных вариантов сетевого графика с расчетами всех его временных параметров и коэффициентов напряженности работ. Процесс «проигрывания» продолжается до тех пор, пока не будет получен приемлемый вариант плана или пока не будет установлено, что все имеющиеся возможности улучшения плана исчерпаны и поставленные перед разработчиком проекта условия невыполнимы.

На практике, кроме соблюдения директивных сроков выполнения комплекса работ, необходимо учитывать вопросы стоимости разработки проекта. Так как при попытках эффективного улучшения составленного плана неизбежно введение дополнительно к оценкам сроков фактора стоимости работ.

ГЛАВА 3. ОПТИМИЗАЦИЯ НА ГРАФАХ

3.1. Задача минимизации дерева расстояний

Имеется n городов $A_1, A_2, \dots, A_n$, которые нужно соединить между собой сетью дорог. Известна стоимость c_{ij} сооружения каждой дороги $A_i A_j$. Какой должна быть сеть дорог, связывающая все города, чтобы стоимость ее сооружения была минимальна?

Алгоритм

1. Выбираем любой узел сети и находим ребро с минимальной величиной c_{ij} . Соединяем два узла (i, j) ребром.
2. Выбираем следующий узел с минимальным значением c_{ij} до уже выбранных узлов. В случае равенства значений c_{ij} выбираем произвольно один из узлов.
3. Если все узлы сети соединены ребрами, то задача решена. В противном случае переходим к шагу 1.

Пример 3.1. Пусть необходимо соединить города А, В, С, Д сетью дорог минимальной стоимости, если известна стоимость сооружения каждой дороги.

	А	В	С	Д
А	-	10	48	45
В	10	-	18	21
С	48	18	-	14
Д	45	21	14	-

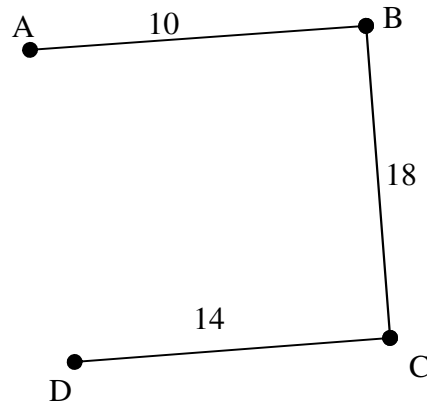


Рис. 3.1.

Решение. В качестве начального узла выбираем узел А. Дорога с минимальной стоимостью связывает узел А с узлом В ($c=10$). Рассматриваем узла А и В. Из них выходят дороги АС ($c=48$), АД ($c=45$), ВС ($c=18$) и ВД ($c=21$). Дорога минимальной стоимости 18 есть дорога ВС. Присоединяем узел С к узлам А и В. Осталось присоединить узел Д. Дорога с минимальной стоимостью $c=14$ есть дорога СД. Таким образом, соединили все узлы сети дорогами (рис. 3.1). При этом минимальная стоимость составит $\min C=10+18+14=42$.

3.2. Поиск кратчайшего пути

Задачи, приводящие к поиску кратчайшего пути.

Пример 3.2. Предположим, что вы хотите проехать из Бостона в Лос-Анджелес, используя магистральные шоссе, соединяющие различные штаты. Как выбрать кратчайший маршрут?

Постройте граф, вершины которого соответствуют точкам соединения рассматриваемых дорог. Пусть каждая дуга графа соответствует шоссе, соединяющей пункты, представленные соответствующими вершинами. Пусть длина дуги определяется длиной (в километрах) соответствующего участка дороги. Теперь задача поиска оптимального маршрута

рута движения между Бостоном и Лос-Анджелесом может быть сведена к задаче отыскания в построенном графе кратчайшего пути между вершинами, соответствующими Бостону, откуда вы начинаете путешествие, и Лос-Анджелесу.

Пример 3.3. Представим себе следующую ситуацию. Коммивояжер планирует поездку из Бостона в Лос Анджелес. Цель поездки – посещение одного из основных клиентов. Коммивояжер собирается воспользоваться той же системой магистральных шоссежных дорог, которая фигурировала в примере 3.2. Помимо основной цели поездки коммивояжер планирует посещение и других пунктов, расположенных на пути из Бостона в Лос-Анджелес, где размещаются другие его клиенты. При этом коммивояжер приблизительно знает, какую сумму комиссионных он заработает после возможной встречи с каждым из клиентов. Какой маршрут поездки из Бостона в Лос-Анджелес следует выбрать коммивояжеру?

В рассматриваемом случае длиной каждой дуги в графе, представляющем рассматриваемую систему магистральных шоссежных дорог, следует считать ожидаемую «чистую стоимость» проезда (транспортные затраты за вычетом ожидаемой суммы комиссионных) по определенному участку данной системы дорог. Теперь можно сказать, что коммивояжер должен ехать по маршруту, соответствующему кратчайшему пути в построенном графе между вершинами «Бостон» и «Лос-Анджелес».

Обратите внимание на то, что в рассматриваемом примере стоимости, характеризующие дуги, будут отрицательными для тех участков маршрута, на которых коммивояжер ожидает получить прибыль, и положительными для тех участков, где он ожидает столкнуться с материальными потерями. (Заметим, что в примере 1 затраты были неотрицательными.) Далее мы увидим, что в каждой из этих двух ситуаций необходимо использовать различные алгоритмы решения задачи о кратчайшем пути.

Описываемый в данном разделе алгоритм позволяет находить в графе кратчайший путь между двумя выделенными вершинами s и t при положительных длинах дуг. Этот алгоритм, предложенный в 1959 г. Дейкстрой, считается одним из наиболее эффективных алгоритмов решения задачи.

Алгоритм Дейкстры можно рассматривать как процедуру наращивания ориентированного дерева с корнем в вершине s . Когда в этой процедуре наращивания достигается конечная вершина t , процедура останавливается.

Шаг 1. Перед началом выполнения алгоритма все вершины и дуги не окрашены. Каждой вершине в ходе проведения алгоритма присваивается число $l(x)$, равное длине кратчайшего пути из начальной вершины s в данную x , включающего только окрашенные вершины. Будем называть данное число весом вершины.

На первом шаге положим $l(s) = 0$ и $l(x) = \infty$. Обозначим через y последнюю на окрашенном пути вершину, т.е. «висячую вершину» дерева. Окрасим вершину s и положим $y = s$.

Шаг 2. Для каждой неокрашенной вершины x рассчитать вес вершины по формуле:

$$l(x) = l(y) + l(x, y) \quad (3.1)$$

Если окажется, что $l(x) = \infty$ для всех неокрашенных вершин x , то нужно закончить процедуру алгоритма: в исходном графе отсутствуют пути из вершины в неокрашенные вершины. В противном случае окрасить ту из вершин x , для которой величина $l(x)$ оказалась наименьшей. Кроме того, необходимо окрасить дугу, ведущую в выбранную на данном шаге вершину x . Положить $y = x$.

Шаг 3. Если $y = t$, то закончить процедуру: кратчайший путь из вершины s в вершину t найден. В противном случае перейти к шагу 2 и провести расчет неокрашенных вершин по формуле (6.1) для всех висячих вершин.

Отметим, что каждый раз, когда окрашивается некоторая вершина (не считая вершины s), окрашивается и некоторая дуга, заходящая в данную вершину. Таким образом, на любом этапе алгоритма в каждую вершину заходит не более чем одна окрашенная дуга. Кроме того, окрашенные дуги не могут образовать в исходном графе цикл, так как в алгоритме не может окрашиваться дуга, концевые вершины которой уже окрашены. Следовательно, можно сделать вывод о том, что окрашенные дуги образуют в исходном графе ориентированное дерево с корнем в вершине s . Это дерево называется *ориентированным деревом кратчайших путей*. Единственный путь от вершины s до любой вершины x , принадлежащей дереву кратчайших путей, является кратчайшим путем между указанными вершинами.

Если кратчайшему пути из вершины s в вершину x в дереве кратчайших путей принадлежит вершина y , то часть этого пути, заключенная между x и y , является кратчайшим путем между этими вершинами. Действительно, если бы между x и y существовал более короткий путь, то упомянутый выше путь между вершинами s и x не мог бы быть кратчайшим.

Поскольку на всех этапах алгоритма Дейкстры окрашенные дуги образуют в исходном графе ориентированное дерево, алгоритм можно рассматривать как процедуру наращивания ориентированного дерева с корнем в вершине s . Когда в этой процедуре наращивания достигается вершина t , процедура может быть остановлена.

Если бы мы хотели определить кратчайшие пути из вершины s во все вершины исходного графа, то процедуру наращивания дерева следовало бы продолжить до тех пор, пока все вершины графа не были бы включены в ориентированное дерево кратчайших путей. При этом для исходного графа было бы получено покрывающее ориентированное дерево (при условии, что в этом графе содержится хотя бы одно такое дерево). Итак, для того, чтобы описанный выше алгоритм позволял получать дерево

кратчайших путей от вершины s до всех остальных вершин, его третий шаг должен быть скорректирован следующим образом: если все вершины оказываются окрашенными, закончить процедуру, в противном случае перейти к шагу 2.

Пример 3.4. Найти кратчайшее расстояние между вершинами s и t графа изображенного на рис. 3.2 с помощью алгоритма Дейкстры.

Решение

Для наглядности проведения процедуры алгоритма Дейкстры составим таблицу 3.1.

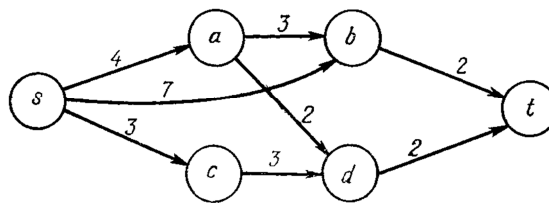


Рис. 3.2.

Таблица 3.1.

№ итерации	Окрашенная вершина (ее вес) / путь	Неокрашенные вершины	Расстояние от окрашенной до неокрашенной вершины,	Путь	Минимальное расстояние	Окрашиваемая вершина / путь
I			4		3	
			7			
			3			
II			4		4	
			7			
			6			
III			7		6	
			6			
			7			
			6			
IV			7		7	
			8			
V			8		8	
			9			

Таким образом, кратчайшее расстояние равно 8 и представляет собой путь .

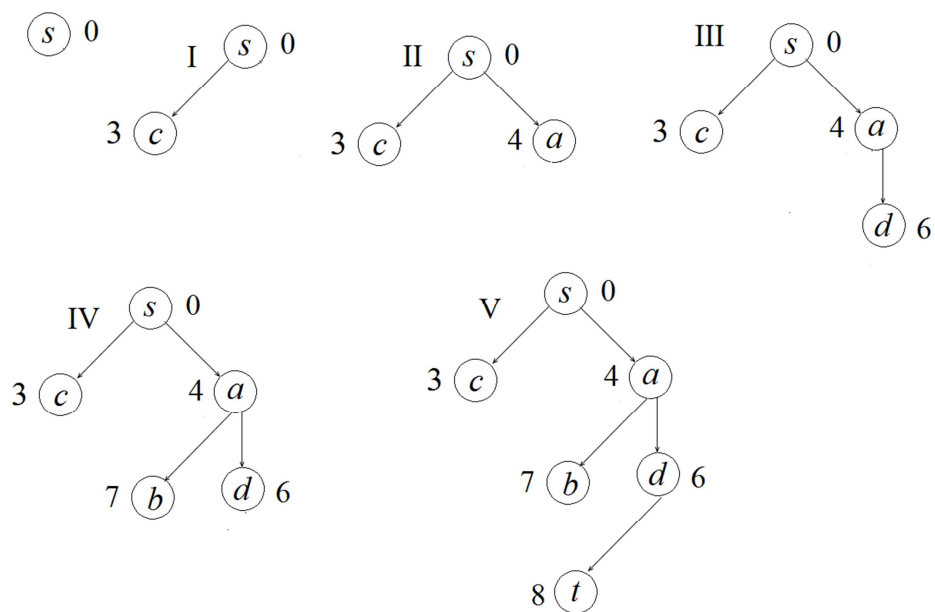


Рис. 3.3. Растущее ориентированное дерево кратчайших путей

ГЛАВА 4. ДИНАМИЧЕСКОЕ ПРОГРАММИРОВАНИЕ

Определение 4.1. Динамическое программирование – это область математики (составная часть математического программирования), разрабатывающая теорию и численные методы нахождения оптимальных планов задач, которые могут быть представлены в виде многошагового процесса.

Динамическое программирование сформировалось в 1950–1953 гг., существенный вклад в его зарождение и становление внес американский математик Р. Беллман. Все задачи линейного программирования обычно решаются в один этап и получили название одноэтапных, одношаговых. Задачи динамического программирования являются многоэтапными, многошаговыми. Собственно, на каждом этапе, за каждым шагом решается частная задача, обусловленная исходной, более общей задачей. В этом смысле динамическое программирование является методом оптимизации многошаговых процессов принятия решений.

Сущность метода динамического программирования заключается в замене одной задачи со многими переменными рядом последовательно решаемых задач с существенно меньшим числом переменных. Оптимизация многошагового процесса решения задачи проводится на основе **принципа оптимальности Р. Беллмана**: оптимальное поведение (оптимальная стратегия) обладает тем свойством, что каково бы не было первоначальное состояние и первоначальное решение, последующее решение должно определять оптимальное поведение относительно состояния, полученного в результате первоначального решения. Суть этого принципа состоит в том, что поэтапный процесс решения должен строиться не на выгоде, получаемой именно на данном шаге, а на общей выгоде от всего процесса.

Принцип Р. Беллмана можно выразить, рассуждая от противного: если не использовать наилучшим образом то, чем мы располагаем сейчас, то

и в дальнейшем не удастся наилучшим образом распорядиться тем, что мы могли бы иметь. В динамическом программировании исследуется особый вычислительный метод, позволяющий осуществлять оптимальное планирование управляемых процессов. Из принципа оптимальности вытекают и правила доминирования, на основе которых на каждом шаге производится сравнение вариантов будущего развития и заблаговременное отсеивание заведомо бесперспективных вариантов. Когда эти правила обращаются в формулы, их называют разрешающими правилами.

Дадим общую постановку задач динамического программирования. Рассматривается управляемый процесс, например, экономический процесс распределения средств между предприятиями, использования ресурсов в течение ряда лет, замены оборудования, пополнения запасов и т.п. В результате управления система (объект управления) S переводится из начального состояния s_0 в состояние s_n . Предположим, что управление можно разбить на n шагов, т.е. решение принимается последовательно на каждом шаге, а управление, переводящее систему S из начального состояния в конечное, представляет собой совокупность n пошаговых управлений.

Обозначим через x_k управление на k -м шаге ($k=1, 2, \dots, n$). Переменные x_k удовлетворяют некоторым ограничениям и в этом смысле называются допустимыми (x_k может быть числом, точкой в n -мерном пространстве, качественным признаком).

Пусть x_k – управление, переводящее систему S из состояния s_{k-1} в состояние s_k . Обозначим через s_k состояние системы после k -го шага управления. Получаем последовательность состояний $s_0, s_1, \dots, s_k, \dots, s_{n-1}, s_n$, которую изобразим кружками (рис. 4.1).

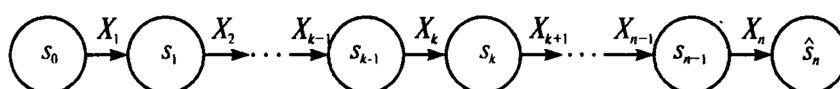


Рис. 4.1.

Показатель эффективности рассматриваемой управляемой системой операции – целевая функция – зависит от начального состояния и управления:

$$Z = F(s_0, X) \quad (4.1)$$

Сделаем несколько предположений.

1. Состояние s_k системы в конце k -го шага зависит только от предшествующего состояния s_{k-1} и управления на k -м шаге X_k (и не зависит от предшествующих состояний и управлений). Это требование называется "отсутствием последствия". Сформулированное положение записывается в виде уравнений

$$s_k = \varphi_k(s_{k-1}, X_k), k = 1, 2, \dots, n,$$

которые называются уравнениями состояний.

2. Целевая функция (4.1) является аддитивной от показателя эффективности каждого шага. Обозначим показатель эффективности k -го шага через

$$Z_k = f_k(s_{k-1}, X_k), k = 1, 2, \dots, n$$

тогда:

$$Z = \sum_{k=1}^n f_k(s_{k-1}, X_k) \quad (4.2)$$

Задача пошаговой оптимизации (задача ДП) формулируется так: определить такое допустимое управление X , переводящее систему S из состояния s_0 в состояние $\hat{s}_0$, при котором целевая функция (4.2) принимает наибольшее (наименьшее) значение.

Выделим особенности модели ДП:

1. Задача оптимизации интерпретируется как пошаговый процесс управления.

2. Целевая функция равна сумме целевых функций каждого шага.

3. Выбор управления на k -м шаге зависит только от состояния системы к этому шагу, не влияет на предшествующие шаги (нет обратной связи).

4. Состояние s_k после k -го шага управления зависит только от предшествующего состояния s_{k-1} и управления X_k (отсутствие последействия).

5. На каждом шаге управление X_k зависит от конечного числа управляющих переменных, а состояние s_k — от конечного числа параметров.

ГЛАВА 5. СИСТЕМЫ МАССОВОГО ОБСЛУЖИВАНИЯ

5.1. Основные положения теории систем массового обслуживания

При исследовании операций часто приходится сталкиваться с системами, предназначенными для многоразового использования при решении однотипных задач. Возникающие при этом процессы получили название процессов обслуживания, а системы – системы массового обслуживания (СМО). К СМО относятся, к примеру, обслуживание покупателей в сфере розничной торговли, медицинское обслуживание населения, ремонт компьютеров и т. д.

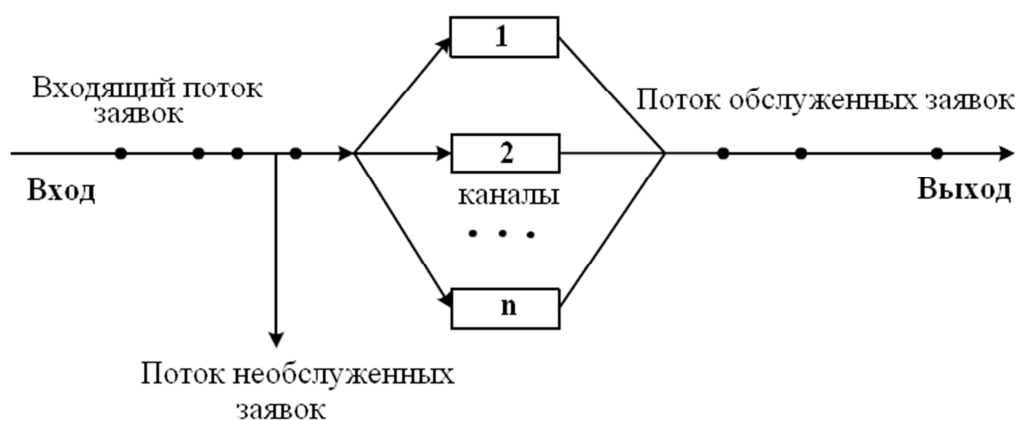


Рис. 5.1. Схема системы массового обслуживания

Заявки в систему массового обслуживания поступают не регулярно, а случайно. При этом они образуют случайный поток заявок (требований). Обслуживание заявки происходит в течение промежутка времени, длительность которого зависит от сложности заявки. Данный промежуток времени также является случайным. Поэтому главной особенностью процессов СМО является случайность. Эти два случайных процесса образуют систему массового обслуживания.

Определение 5.1. Каналами обслуживания СМО называются обслуживающие единицы, например, пункты, приборы и т.п.

Определение 5.2. Системы массового обслуживания – это такие системы, в которые в случайные моменты времени поступают заявки на обслуживание. Поступившие заявки обслуживаются с помощью имеющихся в распоряжении системы каналов обслуживания.

На рисунке 5.1 схематично изображена система массового обслуживания.

Профессиональный менеджер, принимая решение о совершенствовании системы массового обслуживания, оценивает изменения, возникающие в затратах на функционирование системы и в издержках, связанных с ожиданием клиентов. Можно нанять большое количество сотрудников, которые будут быстро обслуживать клиентов. Так, администратор супермаркета может уменьшить очереди в кассы, увеличивая в часы пик количество продавцов и кассиров. Для работы в кассах банков или аэропортов в часы пик могут быть привлечены дополнительные сотрудники. Однако снижение времени ожидания обычно сопряжено с издержками на создание и оснащение рабочих мест, с оплатой труда дополнительного персонала.

Можно экономить на трудозатратах. Но тогда клиент может не дожидаться обслуживания или потерять охоту вернуться еще раз. В последнем случае система массового обслуживания будет нести потери, которые можно назвать издержками ожидания. В некоторых системах обслуживания, например, в скорой помощи, затраты, связанные с длительным ожиданием, могут оказаться чрезвычайно высокими. Основной экономический принцип совершенствования систем массового обслуживания состоит в оценке общих ожидаемых затрат, включающих затраты на обслуживание и потери, которые несет система в результате ожидания клиента.

Таким образом, цель теории массового обслуживания – выработка рекомендаций по рациональному построению СМО, рациональной организации их работы и регулированию потока заявок для обеспечения высокой эффективности функционирования СМО. Для достижения этой цели ста-

вятся задачи теории массового обслуживания, состоящие в установлении зависимостей эффективности функционирования СМО от ее организации (параметров): характера потока заявок, числа каналов и их производительности и правил работы СМО.

Необходимо отметить, что большинство моделей очередей основывается на предположении, что поведение клиентов является стандартным, т.е. каждая поступающая в систему заявка встает в очередь, дожидается обслуживания и не покидает систему до тех пор, пока ее не обслужат. Другими словами, клиент (человек или машина), вставший в очередь, ждет до тех пор, пока он не будет обслужен, не покидает очередь и не переходит из одной очереди в другую.

Жизнь значительно сложнее. На практике клиенты могут покинуть очередь потому, что она оказалась слишком длинной. Может возникнуть и другая ситуация: клиенты ждут своей очереди, но по каким-то причинам уходят необслуженными. Эти случаи также являются предметом теории массового обслуживания, однако здесь не рассматриваются.

По числу каналов СМО делятся на одноканальные и многоканальные. СМО делятся также на два основных вида: СМО с отказами и СМО с очередью.

СМО называется СМО с очередью (ожиданием), если заявка, пришедшая в момент, когда все каналы заняты, не уходит, а становится в очередь на обслуживание.

СМО, которые отказывают в выполнении заявки при занятости всех каналов, называются СМО с отказами.

Системы массового обслуживания, допускающие очередь, но с ограниченным числом мест в ней, называются системами с ограниченной длиной очереди.

Системы массового обслуживания, допускающие очередь, но с ограниченным сроком пребывания каждого требования в ней, называются системами с ограниченным временем ожидания.

Способ отбора для обслуживания заявок называется дисциплиной обслуживания. Различают следующие виды дисциплин:

- Первый пришел — первый обслужился.
- Последний пришел — первый обслужился.
- Обслуживание с ограниченным временем пребывания в очереди.
- Обслуживание приоритетами, когда в первую очередь обслуживаются наиболее важные заявки.

Показатели эффективности СМО описывают ее возможность справляться с потоком заявок.

СМО с отказами.

1. Абсолютная пропускная способность СМО (среднее число заявок, обслуживаемых в единицу времени).

2. Относительная пропускная способность СМО (средняя доля поступивших заявок, обслуживаемых СМО).

3. Вероятность отказа (вероятность того, что заявка покинет СМО необслуженной).

4. Среднее число занятых каналов.

СМО с очередью.

1. Среднее время нахождения заявки в системе.

2. Среднее время нахождения заявки в очереди.

3. Среднее число заявок в очереди.

4. Среднее число заявок в системе.

5. Вероятность того, что канал занят.

Среди экономических характеристик СМО наибольший интерес представляют следующие:

1. Издержки ожидания в очереди;

2. Издержки ожидания в системе;
3. Издержки обслуживания.

Определение 5.3. СМО, у которых входной поток заявок не зависит от того, сколько в данный момент в системе находится количество заявок, называются *разомкнутыми*.

Если характеристики входного потока заявок зависят от того, сколько заявок находится в системе в данный момент, то такие СМО называются *замкнутыми*.

Примером замкнутой системы может служить ремонтная мастерская таксомоторного парка с заданным числом машин. Чем больше машин находится в системе, тем меньше их эксплуатируют и тем меньше становится интенсивность потока вновь поступающих на ремонт машин.

По количеству этапов обслуживания СМО делятся на однофазные и многофазные системы. Если каналы СМО однородны, т.е. выполняют одну и ту же операцию обслуживания, то такие СМО называются *однофазными*. Если каналы обслуживания расположены последовательно, и они неоднородны, так как выполняют различные операции обслуживания (т.е. обслуживание состоит из нескольких последовательных этапов или фаз), то СМО называется *многофазной*. Примером работы многофазной СМО является обслуживание автомобилей на станции технического обслуживания (мойка, диагностирование и т.д.).

Работа СМО является случайным процессом с дискретными состояниями и непрерывным временем.

Определение 5.4. *Случайным процессом* называется соответствие, при котором каждому значению аргумента (в данном случае – моменту из промежутка времени проводимого опыта) ставится в соответствие случайная величина (в данном случае – состояние СМО).

Определение 5.5. Случайный процесс называется *марковским*, или *случайным процессом без последствий*, если для любого момента времени

вероятность любого состояния системы в будущем зависит только от ее состояния в настоящем и не зависит от того, как система пришла в это состояние.

Примером марковского процесса являются показания счетчика в такси. Показания счетчика в будущем зависят от показаний в настоящий момент и не зависят от того, в какие моменты времени изменялись показания счетчика до настоящего момента.

На практике марковские процессы в чистом виде обычно не встречаются, но нередко приходится иметь дело с процессами, которые можно приближенно считать марковскими, т.е. такие, для которых влиянием «предыстории» можно пренебречь. Кроме того, существуют приемы, позволяющие сводить немарковские случайные процессы к марковским. Например, можно вводить в состав параметров, характеризующих настоящее состояние системы, те параметры из прошлого, от которых зависит будущее (в этом случае говорят о «марковизации» случайного процесса). Правда, такая процедура нередко приводит к сильному усложнению математического аппарата. Существуют и другие приемы сведения немарковских случайных процессов к марковским.

В большинстве случаев для обоснованных рекомендаций по практическому управлению СМО совсем не требуется знаний точных ее характеристик, вполне достаточно иметь их приближенные значения.

Определение 5.6. *Потоком событий* называется последовательность однородных событий, следующих одно за другим в случайные моменты времени.

Определение 5.7. *Интенсивностью* потока событий называется среднее число событий, поступающих в систему за единицу времени.

В СМО требования могут поступать через различные промежутки времени, т.е. интервал между соседними поступающими требованиями – величина случайная.

Определение 5.8. *Стационарным* называется поток, у которого вероятностные характеристики не зависят от времени. В частности, интенсивность потока есть величина постоянная.

Определение 5.9. *Потоком событий без последствий* называется поток, число событий которого, попавших на заданный временной интервал, не зависит от числа событий, попавших на другие интервалы.

Определение 5.10. *Ординарным потоком событий* называется поток, в котором вероятность попадания на элементарный временной интервал двух и более событий пренебрежимо мала по сравнению с вероятностью попадания одного события.

В настоящее время теоретически наиболее разработаны методы решения таких задач СМО, в которых входящий поток требований является *простейшим (пуассоновским)*.

Определение 5.11. *Простейшим потоком* событий называется поток, если он является стационарным, ординарным и не имеет последствий.

Для простейшего потока частота поступления требований в систему подчиняется закону Пуассона, т.е. вероятность поступления за время τ ровно m требований задается формулой:

$$P_m(\tau) = \frac{(\lambda\tau)^m}{m!} e^{-\lambda\tau} \quad (5.1)$$

λ – интенсивность потока событий, поступающих в СМО. Для закона Пуассона математическое ожидание равно:

$$a = \sigma = \frac{1}{\lambda} \quad (5.2)$$

Вероятность того, что на временном интервале τ появится хотя бы одно событие, определяется соотношением:

$$P(\tau) = 1 - P_0 = 1 - e^{-\lambda\tau} \quad (5.3)$$

где $P_0 = e^{-\lambda\tau}$ – вероятность того, что за интервал τ не произойдет ни одного события.

Вероятность попадания на элементарный временной интервал $\Delta t = \tau$ хотя бы одного события потока находим из соотношения (5.3), заменив функцию $e^{-\lambda\tau}$ двумя первыми членами ее разложения в ряд Маклорена по степеням Δt :

$$P(\Delta t) = 1 - e^{-\lambda\Delta t} \approx \lambda \cdot \Delta t \quad (5.4)$$

Следует отметить, что на практике поток заявок, разумеется, не всегда подчиняется пуассоновскому распределению (он может иметь какое-то другое распределение). Поэтому требуется проводить предварительные исследования для того, чтобы проверить, что пуассоновское распределение может служить хорошей аппроксимацией.

Если пуассоновский поток нестационарный, т.е. $\lambda = \lambda(t)$, то для расчетов либо весь временной интервал $[0, T)$ разбивается на n временных отрезков, вообще говоря, не равных, и для каждого из них находятся $\lambda_i(t) = \lambda_i \cdot t$, где $t \in [t_{i-1}, t_i)$, $\sum_{i=1}^n [t_{i-1}, t_i) = T$, либо используются численные методы. В зависимости от поставленной цели выбирается аналитический подход либо численный.

Случайный процесс, протекающий в СМО, состоит в том, что система в случайные моменты времени переходит из одного состояния в другое: меняется число занятых каналов, число заявок, стоящих в очереди, и т.п. Это означает, что СМО представляет собой физическую систему дискретного типа с конечным (или счетным) множеством состояний, а переход системы из одного состояния в другое происходит скачком, в момент, когда осуществляется какое-то событие (приход новой заявки, освобождение канала, уход заявки из очереди и т.п.).

Случайные процессы с дискретными состояниями (не более чем счетным множеством состояний) бывают двух типов: с дискретным или непрерывным временем. Первые отличаются тем, что переходы из состояния в состояние могут происходить только в строго определенные, разделенные

конечными интервалами моменты времени. Случайные процессы с непрерывным временем отличаются тем, что переход системы из состояния в состояние возможен в любой момент времени.

При анализе случайных процессов с дискретными состояниями пользуются графом состояний, где прямоугольниками изображаются состояния, переходы из состояния в состояние – стрелками. Если у стрелок представлены интенсивности, то граф состояний называется размеченным.

Рассмотрим математическое описание марковского случайного процесса с дискретными состояниями и непрерывным временем на примере следующего случайного процесса. Устройство S состоит из двух узлов, каждый из которых может выйти из строя в любой случайный момент времени, после чего мгновенно начинается его ремонт, продолжающийся заранее неизвестное время.

Возможные состояния системы: S_0 – оба узла исправны; S_1 – первый узел ремонтируется, второй исправен; S_2 – второй узел ремонтируется, первый исправен; S_3 – оба узла ремонтируются. Граф системы приведен на рисунке 5.3. Будем полагать, что все переходы системы из состояния в состояние происходят под воздействием простейших потоков событий с интенсивностями λ_{ij} ($i, j = 0, 1, 2, 3$).

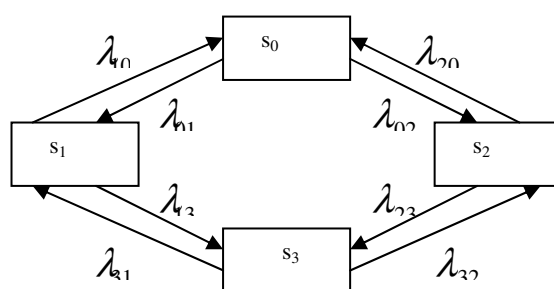


Рис. 5.2

Колмогоровым были получены уравнения, дающие возможность найти все вероятности состояний как функции времени. Особый интерес

среди них представляют вероятности $P_i(t)$ при $t \rightarrow \infty$, называемые *предельными вероятностями*.

Определение 5.12. Предельные вероятности состояний S_i показывают среднее относительное время пребывания системы в этом состоянии.

Система уравнений Колмогорова может быть получена непосредственно из графа состояний по следующему правилу.

Слева в уравнениях стоит предельная вероятность данного состояния p_i , умноженная на суммарную интенсивность всех потоков, ведущих из данного состояния, а справа – сумма произведений интенсивностей всех потоков, входящих в i -е состояние, на вероятности тех состояний, из которых эти потоки исходят. Для нашего случая:

$$\begin{cases} (\lambda_{01} + \lambda_{02}) p_0 = \lambda_{10} p_1 + \lambda_{20} p_2, \\ (\lambda_{10} + \lambda_{13}) p_1 = \lambda_{01} p_0 + \lambda_{31} p_3, \\ (\lambda_{20} + \lambda_{23}) p_2 = \lambda_{02} p_0 + \lambda_{32} p_3, \\ (\lambda_{31} + \lambda_{32}) p_3 = \lambda_{13} p_1 + \lambda_{23} p_2 \end{cases} \quad (5.5)$$

Пример 5.1. Найти предельные вероятности для системы, граф состояний которой изображен на рисунке 5.3, если $\lambda_{01} = 1$, $\lambda_{02} = 2$, $\lambda_{10} = 3$, $\lambda_{13} = 2$, $\lambda_{20} = 3$, $\lambda_{23} = 1$, $\lambda_{31} = 3$, $\lambda_{32} = 1$.

Решение:

Система алгебраических уравнений, описывающих стационарный режим, для данной системы имеет вид:

$$\begin{cases} 3p_0 = 3p_1 + 3p_2 \\ 5p_1 = p_0 + 3p_3 \\ 4p_2 = 2p_0 + p_3 \\ p_0 + p_1 + p_2 + p_3 = 1 \end{cases}$$

Последнее уравнение из данной системы является условием нормировки. Решение системы дает: $p_0=0,425$, $p_1 = 0,175$, $p_2 = 0,25$, $p_3 = 0,15$, т.е. в предельном стационарном режиме система S в среднем 42,5% времени будет находиться в состоянии S_0 , 17,5% – в состоянии S_1 , 25 % – в состоянии S_2 и 12,5% в состоянии S_3 .

В теории массового обслуживания широко распространен специальный класс случайных процессов – так называемые процессы гибели и размножения. Название это связано с рядом биологических задач, где этот процесс служит математической моделью изменения численности биологических популяций. Граф состояний процесса гибели и размножения имеет вид, изображенный на рисунке 5.3.

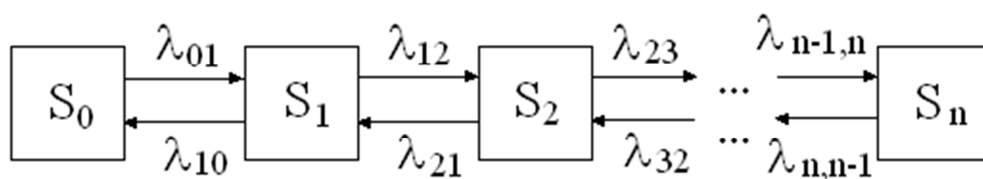


Рис. 5.3. Граф гибели и размножения

Особенностью данного графа является то, что каждое состояния, кроме первого и последнего, связано только с предыдущим и последующим состояниями. Составим и решим алгебраические уравнения для предельных вероятностей состояний.

В соответствии с правилом составления таких уравнений получим для состояния S_0 :

$$\lambda_{01}p_0 = \lambda_{10}p_1 \quad (5.6)$$

Для состояния S_1 :

$$(\lambda_{12} + \lambda_{10})p_1 = \lambda_{01}p_0 + \lambda_{21}p_2,$$

которое с учетом (5.6) приводится к виду

$$\lambda_{12} p_1 = \lambda_{21} p_2$$

Аналогично, записывая уравнение для предельных вероятностей других состояний, получим следующую систему линейных уравнений:

$$\begin{cases} \lambda_{01} p_0 = \lambda_{10} p_1, \\ \lambda_{12} p_1 = \lambda_{21} p_2, \\ \dots\dots\dots \\ \lambda_{n-1,n} p_{n-1} = \lambda_{n,n-1} p_n \end{cases} \quad (5.7)$$

Используя нормировочное условие:

$$p_0 + p_1 + p_2 + \dots + p_n = 1 \quad (5.8)$$

Из первого уравнения выразим p_1 через p_0 :

$$p_1 = \frac{\lambda_{01}}{\lambda_{10}} p_0 \quad (5.9)$$

Из второго уравнения выразим p_2 через p_0 , используя подстановку (5.9)

$$p_2 = \frac{\lambda_{12}}{\lambda_{21}} p_1 = \frac{\lambda_{12}\lambda_{01}}{\lambda_{21}\lambda_{10}} p_0 \quad (5.10)$$

Из третьего уравнения выразим p_3 через p_0 , используя (5.10):

$$p_3 = \frac{\lambda_{13}}{\lambda_{31}} p_2 = \frac{\lambda_{12} \lambda_{01} \lambda_{13}}{\lambda_{21} \lambda_{10} \lambda_{31}} p_0$$

И т.д.

$$p_n = \frac{\lambda_{n-1,n} \dots \lambda_{12} \lambda_{01}}{\lambda_{n,n-1} \dots \lambda_{21} \lambda_{10}} p_0 \quad (5.11)$$

В числителе формулы (5.11) стоит произведение всех интенсивностей, стоящих у стрелок, ведущих слева направо до данного состояния S_k ($k = 1, 2, \dots, n$), а в знаменателе – произведения всех интенсивностей, стоящих у стрелок, ведущих справа налево из состояния S_k (рис. 5.4).

Преобразуем формулу (5.11), используя условие (5.8), к виду:

$$p_0 = \left(1 + \frac{\lambda_{01}}{\lambda_{10}} + \frac{\lambda_{12}\lambda_{01}}{\lambda_{21}\lambda_{10}} + \dots + \frac{\lambda_{n-1,n}\dots\lambda_{12}\lambda_{01}}{\lambda_{n,n-1}\dots\lambda_{21}\lambda_{10}} \right)^{-1} \quad (5.12)$$

5.2. Системы массового обслуживания с отказами

Рассмотрим многоканальные системы массового обслуживания с отказами. Анализ такой СМО строится на базе классической задачи Эрланга (датского инженера, основателя теории СМО).

Имеется n каналов, на которые поступает поток заявок с интенсивностью λ . Поток обслуживаний имеет интенсивность μ . Требуется определить предельные вероятности состояний системы и показатели ее эффективности.

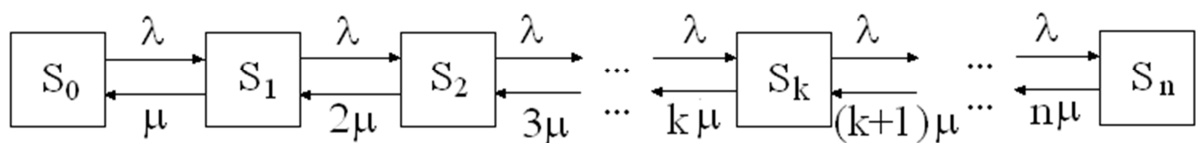


Рис. 5.4. Граф состояний для многоканальной СМО с отказами

Система S (СМО) имеет следующие состояния: S_0 (все каналы свободны), S_1 (один канал занят), S_2 (два канала заняты), ..., S_k (k каналов занято), ..., S_n (все каналы заняты). Принцип работы такой многоканальной СМО с отказами рассмотрим с помощью графа гибели и размножения, представленного для данного случая на рисунке 5.4.

Поток заявок последовательно переводит систему из любого состояния в соседнее правое с одной и той же интенсивностью λ . Интенсивность же потока обслуживаний, переводящих систему из любого правого состояния в соседнее левое, постоянно меняется в зависимости от состояния. Действительно, если СМО находится в состоянии S_2 (два канала заняты), то она может перейти в состояние S_1 (один канал занят), когда закончит обслуживание либо первый, либо второй канал, т.е. суммарная интенсивность их потоков обслуживаний будет 2μ и т.д.

Используя формулу (5.12) для схемы гибели и размножения, получим для предельной вероятности состояния S_0

$$p_0 = \left(1 + \frac{\lambda}{\mu} + \frac{\lambda^2}{2!\mu^2} + \dots + \frac{\lambda^n}{n!\mu^n} \right)^{-1} \quad (5.13)$$

где члены разложения $\frac{\lambda}{\mu}, \frac{\lambda^2}{2!\mu^2}, \dots, \frac{\lambda^n}{n!\mu^n}$ будут представлять собой коэффициенты при p_0 в выражениях для предельных вероятностей $p_1, p_2, \dots, p_n$.

Величину $\rho = \frac{\lambda}{\mu}$ называют *интенсивностью нагрузки системы*. Она выражает среднее число заявок, приходящее за среднее обслуживания одной заявки. Запишем формулу 5.13 с помощью показателя ρ :

$$p_0 = \left(1 + \rho + \frac{\rho^2}{2!} + \dots + \frac{\rho^n}{n!} \right)^{-1} \quad (5.14)$$

$$p_1 = \rho \cdot p_0, \quad p_2 = \frac{\rho^2}{2!} p_0, \dots, \quad p_n = \frac{\rho^n}{n!} p_0$$

Формулы (5.14) получили название формул Эрланга в честь основателя теории массового обслуживания.

Показатели эффективности многоканальной СМО с отказами:

1) Доля потерянных требований или вероятность отказа СМО есть предельная вероятность того, что все n каналов системы будут заняты:

$$p_{\text{отк}} = \frac{\rho^n}{n!} p_0 \quad (5.15)$$

2) Относительная пропускная способность – вероятность того, что заявка будет обслужена:

$$Q = 1 - p_{\text{отк}} = 1 - \frac{\rho^n}{n!} p_0 \quad (5.16)$$

3) Абсолютная пропускная способность:

$$A = \lambda Q = \lambda \left(1 - \frac{\rho^n}{n!} p_0 \right) \quad (5.17)$$

4) Среднее число занятых каналов с учетом того, что каждый занятый канал обслуживает в среднем μ заявок в единицу времени равно:

$$\bar{k} = \frac{A}{\mu} = \rho(1 - p_{\text{отк}}) \quad (5.18)$$

Пример 5.2. В магазине принимают заказы по двум телефонам. Среднее число поступающих в течение часа заказов – 70, а среднее время оформления заказа – 3 мин. Определить показатели эффективности СМО.

Решение:

Из условия следует, что интенсивность потока обслуживания $\mu = 20$ заявок в час. Параметр загрузки системы: $\rho = \frac{\lambda}{\mu} = \frac{70}{20} = 3,5$.

Вероятность того, что все каналы будут свободны:

$$p_0 = \left(1 + \rho + \frac{\rho^2}{2!}\right)^{-1} = \left(1 + 3,5 + \frac{3,5^2}{1 \cdot 2}\right)^{-1} \approx 0,094$$

Доля потерянных заказов:

$$p_{\text{отк}} = \frac{\rho^2}{2!} p_0 = \frac{3,5^2}{1 \cdot 2} \cdot 0,094 \approx 0,576$$

Пропускная способность системы:

$$Q = 1 - 0,576 = 0,424$$

Абсолютная пропускная способность:

$$A = 70 \cdot 0,424 \approx 29,7$$

Среднее число занятых каналов:

$$\bar{k} = \frac{29,7}{20} \approx 1,48$$

Оба телефона заняты, и заказ невозможно сделать в 57,6 % случаев.

Показатель $P_{\text{отк}} = \frac{\rho^n}{n!} p_0$ означает долю времени, когда в системе заняты

все n каналов.

Рассмотрим одноканальную СМО с отказами. Имеется один канал, на который поступают заявки (поток) с интенсивностью λ . Поток обслуживаний имеет интенсивность μ .

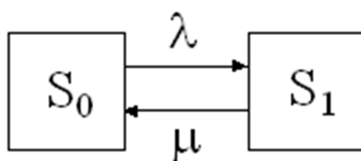


Рис. 5.5. Граф гибели и размножения для одноканальной СМО с отказами

На рисунке 5.5. изображен граф гибели и размножения для данной СМО. Система имеет два состояния: S_0 – канал свободен, S_1 – канал занят.

Найти предельные вероятности состояний системы и показатели ее эффективности. Составляем уравнения Колмогорова

$$\begin{cases} \lambda \cdot p_0 = \mu \cdot p_1 \\ \mu \cdot p_1 = \lambda \cdot p_0 \end{cases} \quad (5.19)$$

т.е. система выражается в одно уравнение. С учетом равенства $p_0 + p_1 = 1$, получим предельные вероятности:

$$p_0 = \frac{\mu}{\lambda + \mu}, \quad p_1 = \frac{\lambda}{\lambda + \mu} \quad (5.20)$$

которое выражают среднее относительное время пребывания системы в состоянии S_0 и S_1 .

Показатели эффективности одноканальной СМО с отказами:

1) Относительная пропускную способность системы:

$$Q = \frac{\mu}{\lambda + \mu} \quad (5.21)$$

2) Вероятность отказа $P_{отк}$:

$$P_{отк} = \frac{\lambda}{\lambda + \mu} \quad (5.22)$$

3) Абсолютная пропускная способность СМО:

$$A = \lambda Q = \frac{\lambda \mu}{\lambda + \mu} \quad (5.23)$$

4) Среднее число занятых каналов:

$$\bar{k} = \frac{A}{\mu} = \frac{\lambda}{\lambda + \mu} \quad (5.24)$$

Пример 5.3. В магазине принимают заказы по одному телефону. Среднее число поступающих в течение часа заказов – 70, а среднее время оформления заказа – 2 мин. Определить показатели эффективности СМО.

Решение:

Имеем $\lambda = 70 (1/ч)$, $\bar{t}_{об} = 2 \text{ мин}$.

Интенсивность потока обслуживаний $\mu = \frac{1}{\bar{t}_{об}} = \frac{1}{2} = 30 (1/ч)$

Относительная пропускная способность Q системы:

$$Q = \frac{\mu}{\lambda + \mu} = \frac{30}{70 + 30} = 0,3$$

Доля потерянных заказов:

$$P_{отк} = \frac{\lambda}{\lambda + \mu} = \frac{70}{70 + 30} = 0,7$$

Следовательно, только 30% поступающих заказов принимаются, в 70% случаев этого сделать нельзя, т.к. телефон оказывается занятым.

Абсолютная пропускная способность СМО равна:

$$A = 70 \cdot 0,3 = 21,$$

т.е. в среднем в час будут обслужены 21 заказ. Очевидно, что при наличии одного телефона СМО плохо справляется с потоком заказов.

5.3. СМО с очередью

На практике часто встречаются одноканальные СМО с неограниченной очередью (например, телефон-автомат с одной кабиной, очередь на прием к врачу, очередь на проезд по мосту при движении с одной полосой,

очередь на входе в автобус при наличии устройства автоматизированного контроля проезда пассажиров и т.д.). Итак, имеется одноканальная СМО с очередью, на которую не наложены никакие ограничения (ни по длине очереди, ни по времени ожидания). Поток заявок, поступающих в СМО, имеет интенсивность λ , а поток обслуживаний – интенсивность μ . Необходимо найти предельные вероятности состояний и показатели эффективности СМО.

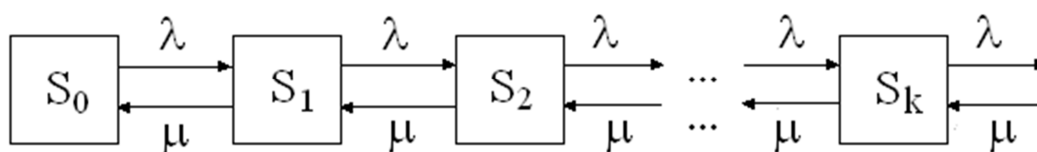


Рис. 5.6. Граф состояний одноканальной СМО с неограниченной очередью

Система может находиться в одном из состояний $S_0, S_1, S_2, \dots, S_k, \dots$ по числу заявок, находящихся в СМО: S_0 – канал свободен, S_1 – канал занят (обслуживает заявку), очереди нет, S_2 – канал занят, одна заявка стоит в очереди, $\dots$, S_k – канал занят, $k-1$ заявок в очереди и т.д.

Граф состояний представлен на рисунке 5.6.

Прежде чем найти выражения для предельных вероятностей, необходимо быть уверенным в их существовании. Если $\rho = \frac{\lambda}{\mu} < 1$, т.е. в единицу времени среднее число пришедших заявок меньше среднего числа обслуженных заявок, то предельные вероятности существуют, если же $\rho = \frac{\lambda}{\mu} > 1$, то очередь растет до бесконечности.

Предельная вероятность p_0 для состояния S_0 имеет вид:

$$p_0 = (1 + \rho + \rho^2 + \dots + \rho^k + \dots)^{-1} \quad (5.25)$$

Так как $\rho < 1$, то в скобках имеет место геометрический ряд, который сходится к сумме, равной $\frac{1}{1-\rho}$, поэтому:

$$p_0 = 1 - \rho \quad (5.26)$$

Предельные вероятности для других состояний системы:

$$p_1 = \rho \cdot p_0, \quad p_2 = \rho^2 \cdot p_0, \dots, \quad p_k = \rho^k \cdot p_0, \dots \quad (5.27)$$

Предельные вероятности с учетом формулы (5.26) приобретают вид:

$$p_1 = \rho \cdot (1 - \rho), \quad p_2 = \rho^2 \cdot (1 - \rho), \dots, \quad p_k = \rho^k \cdot (1 - \rho), \dots \quad (5.28)$$

Предельные вероятности $p_0, p_1, \dots, p_n$ образуют убывающую геометрическую прогрессию со знаменателем $\rho < 1$. Следовательно, вероятность p_0 – наибольшая. Это означает, что если СМО справляется с потоком заявок (при $\rho < 1$), то наиболее вероятным будет отсутствие заявок в системе.

Вычислим среднее число заявок в системе $L_{\text{сист}}$. Так как количество заявок может принимать значения $0, 1, 2, \dots, k, \dots$ то по формуле математического ожидания можно записать:

$$L_{\text{сист}} = 0p_0 + 1p_1 + 2p_2 + 3p_3 + \dots + kp_k = \sum_{k=1}^{\infty} kp_k$$

С учетом формул (5.26):

$$L_{\text{сист}} = (1 - \rho) \sum_{k=1}^{\infty} k\rho^k \quad (5.29)$$

Можно показать, что формула (5.29) преобразуется при $\rho < 1$ к виду:

$$L_{\text{сист}} = \frac{\rho}{1 - \rho} \quad (5.30)$$

Теперь определим среднее число заявок в очереди (длину очереди) $L_{\text{оч}}$. Длина очереди – это разница между общим число заявок в системе и заявками, находящимися на обслуживании:

$$L_{\text{оч}} = L_{\text{сист}} - L_{\text{об}}, \quad (5.31)$$

где $L_{\text{об}}$ – среднее число заявок, ожидающих обслуживания.

Так как рассматриваемая СМО одноканальная, то обслуживаться может только одна заявка, а остальные заявки ждут своей очереди. Среднее число заявок под обслуживанием определим по формуле математического

ожидания числа заявок под обслуживанием, принимающего значение 0 (если канал свободен), либо 1 (если канал занят)

$$L_{об} = 0p_0 + 1(1 - p_0) = 1 - p_0 = 1 - (1 - \rho) = \rho \quad (5.32)$$

Теперь:

$$L_{оч} = L_{сист} - \rho = \frac{\rho}{1 - \rho} - \rho = \frac{\rho^2}{1 - \rho} \quad (5.33)$$

Среднее время пребывания заявки в системе (очереди) равно среднему числу заявок в системе (очереди), деленному на интенсивность потока заявок, т.е.

$$T_{сист} = \frac{1}{\lambda} L_{сист}, \quad T_{оч} = \frac{1}{\lambda} L_{оч} \quad (5.34)$$

Формулы (5.34) называются формулами Литтла. Они вытекают из того, что в предельном, стационарном режиме среднее число заявок, пребывающих в систему, равно среднему числу заявок, покидающих ее: оба потока заявок имеют одну и ту же интенсивность λ .

С учетом (5.32) и (5.34) имеем:

$$T_{сист} = \frac{\rho}{\lambda(1 - \rho)} \quad (5.35)$$

$$T_{оч} = \frac{\rho^2}{\lambda(1 - \rho)} \quad (5.36)$$

Пример 5.4. На оптовую базу поступают на разгрузку пять автомобилей в час ($\lambda = 5$). Среднее время разгрузки одного автомобиля 10 мин. Определить показатели эффективности СМО.

Решение:

Имеем одноканальную СМО с неограниченной очередью. Интенсивность обслуживания автомобилей $\mu = \frac{1}{t_{об}} = \frac{1}{1/6} = 6$

Интенсивность нагрузки канала: $\rho = \frac{\lambda}{\mu} = \frac{5}{6} \approx 0,83 < 1$.

Среднее количество машин, находящихся в системе:

Среднее время обслуживания автомобиля:

Длина очереди (среднее число автомобилей, ожидающих разгрузки):

Среднее время ожидания автомобиля в очереди:

Имеется n -канальная СМО с неограниченной очередью. Поток заявок, поступающих в СМО, имеет интенсивность λ , а поток обслуживаний – интенсивность μ . Необходимо найти предельные вероятности состояний СМО и показатели ее эффективности.

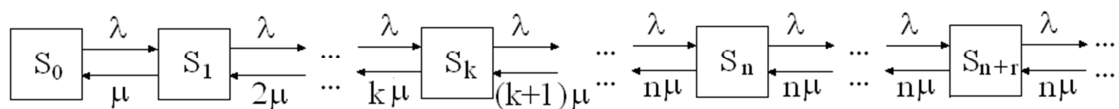


Рис. 5.7. Граф состояний многоканальной СМО с неограниченной очередью

Система может находиться в одном из состояний $S_0, S_1, \dots, S_k, \dots, S_n, S_{n+1}, \dots, S_r, \dots$ по числу заявок, находящихся в СМО: S_0 – канал свободен, S_1 – один канал занят (обслуживает заявку), очереди нет, S_2 – два канала занято, ..., S_k – занято k каналов, остальные свободны; ...; S_n – заняты все n каналов (очереди нет); ...; S_{n+r} – заняты все n каналов, r заявок в очереди, Граф состояний приведен на рисунке 5.7.

В отличие от одноканальной СМО интенсивность потока обслуживаний не остается постоянной, а по мере увеличения числа заявок в СМО от 0 до n увеличивается от величины μ до $n\mu$, т.к. соответственно увеличивается число каналов обслуживания. При числе заявок больше, чем n , интенсивность потока обслуживаний сохраняется равной $n\mu$. Если в системе n

каналов обслуживания с интенсивностью μ , интенсивность входящего потока равна λ , то, чтобы очередь не стала бесконечно большой, необходимо выполнение условия стационарности $\frac{\lambda}{\mu} < n$.

Это условие означает, что суммарная скорость обслуживания всех каналов системы $n\mu$ должна превосходить скорость поступления требований λ , иначе система не справится с обслуживанием потока.

Данное условие характерно только для систем с очередью в отличие от систем с отказом, т.к. все поступившие требования должны получить обслуживание.

Используя формулы (5.13) для процесса гибели и размножения, можно получить формулы для предельных вероятностей состояний n -канальной СМО с неограниченной очередью:

$$p_0 = \left(1 + \rho + \frac{\rho^2}{2!} + \dots + \frac{\rho^n}{n!} + \frac{\rho^{n+1}}{n!(n-\rho)} \right)^{-1} \quad (5.37)$$

$$p_1 = \frac{\rho}{1} \cdot p_0, \dots, p_k = \frac{\rho^k}{k!} p_0, \dots, p_n = \frac{\rho^n}{n!} p_0 \quad (5.38)$$

$$p_{n+1} = \frac{\rho^{n+1}}{n \cdot n!} p_0, \dots, p_{n+r} = \frac{\rho^{n+r}}{n^r \cdot n!} p_0$$

Вероятность того, что в системе заняты обслуживанием все каналы, определяется по формуле:

$$p_{\text{очер}} = \frac{\rho^{n+1}}{(n-\rho)n!} p_0 \quad (5.39)$$

Для n -канальной СМО с неограниченной очередью, используя прежние приемы, можно найти:

1) среднее число занятых каналов:

$$\bar{k} = \rho \quad (5.40)$$

2) среднее число заявок в очереди:

$$L_{\text{оч}} = \frac{\rho^{n+1} p_0}{(n-1)!(n-\rho)^2} \quad (5.41)$$

3) среднее число заявок в системе:

$$L_{\text{сист}} = L_{\text{оч}} + \rho \quad (5.42)$$

4) среднее время обслуживания заявки:

$$T_{\text{сист}} = \frac{L_{\text{сист}}}{\lambda} \quad (5.43)$$

5) среднее время ожидания обслуживания:

$$T_{\text{оч}} = \frac{L_{\text{оч}}}{\lambda} \quad (5.44)$$

Пример 5.5. К двум продавцам поступает на обслуживание поток покупателей с интенсивностью 220 человек в час. Каждый из продавцов затрачивает на обслуживание покупателя в среднем 30 секунд. Определите среднюю длину очереди и показатели занятости продавцов.

Решение:

$$\lambda = 220 \text{ чел/час}, \mu = \frac{1}{30} \cdot 60 \cdot 60 = 120 \text{ чел/час}, n = 2$$

$$\rho = \frac{\lambda}{\mu} = \frac{220}{120} = \frac{11}{6} - \text{интенсивность загрузки}$$

$$\bar{k} = \rho = \frac{11}{6} = 1\frac{5}{6} - \text{среднее число занятых обслуживанием каналов}$$

$$L_{\text{оч2}} = \frac{\rho^3}{4 - \rho^2} = \frac{\left(\frac{11}{6}\right)^3}{4 - \left(\frac{11}{6}\right)^2} \approx 10 \text{ чел} - \text{средняя длина очереди}$$

$$p_0 = \frac{2 - \rho}{2 + \rho} \approx 0,04 - \text{доля времени простоя продавцов}$$

$$p_1 = \rho \cdot p_0 = \frac{11}{6} \cdot \frac{1}{23} \approx 0,08 - \text{доля времени занятости одного из двух про-}$$

давцов

$$p_2 = \frac{\rho^2}{2 + \rho} = \frac{121}{36 \left(2 + \frac{11}{6}\right)} \approx 0,88 - \text{доля времени занятости двух продавцов}$$

Пусть число заявок, находящихся в очереди, не превосходит m . Если новая заявка поступает в момент, когда все места в очереди заняты, то она получит отказ. Сначала рассмотрим одноканальную СМО с ограниченной очередью. На рисунке 5.8 показан граф состояний для такой СМО.

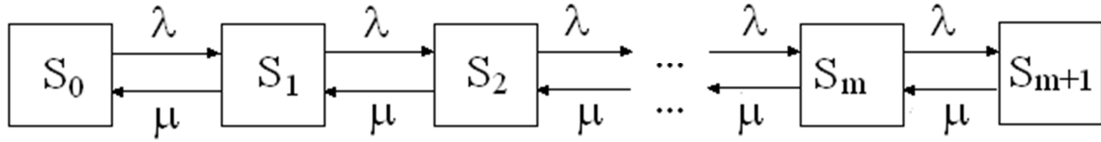


Рис. 5.8. Граф состояний одноканальной СМО с ограниченной очередью

Система может находиться в одном из состояний $S_0, S_1, S_2, \dots, S_m$ по числу заявок, находящихся в СМО: S_0 – канал свободен, S_1 — канал занят (обслуживает заявку), очереди нет, S_2 – канал занят, одна заявка стоит в очереди, ..., S_m – канал занят, $m-1$ заявок в очереди, S_{m+1} , m заявок в очереди.

Показатели эффективности работы одноканальной СМО с ограниченной очередью:

1) Предельные вероятности:

$$P_0 = \frac{1-\rho}{1-\rho^{m+2}}, \dots, P_r = \rho^r P_0, \dots \quad (5.45)$$

2) Вероятность отказа:

$$P_{отк} = P_{m+1} = \rho^{m+1} P_0 \quad (5.46)$$

3) Относительная пропускная способность:

$$Q = 1 - P_{отк} = 1 - \rho^{m+1} P_0 \quad (5.47)$$

4) Абсолютная пропускная способность:

$$A = \lambda Q = \lambda(1 - \rho^{m+1} P_0) \quad (5.48)$$

5) Среднее число заявок в очереди:

$$L_{оч} = \rho^2 \frac{[1 - \rho^m(m+1 - m\rho)]}{(1 - \rho^{m+2})(1 - \rho)} \quad (5.49)$$

6) Среднее число заявок под обслуживанием:

$$L_{обсл} = 1 - P_0 \quad (5.50)$$

7) Среднее число заявок в системе:

$$L_{сист} = L_{обсл} + L_{оч} \quad (5.51)$$

На рисунке 5.9 представлен граф состояний многоканальной СМО с ограниченной очередью.

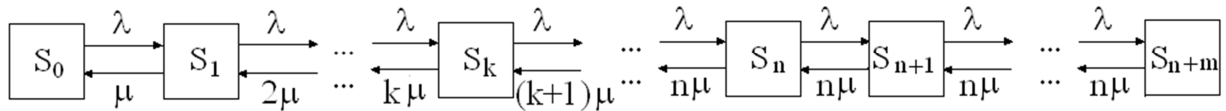


Рис. 5.9. Граф состояний одноканальной СМО с ограниченной очередью

Система может находиться в одном из состояний $S_0, S_1, \dots, S_k, \dots, S_n, S_{n+1}, \dots, S_m$ по числу заявок, находящихся в СМО: S_0 – канал свободен, S_1 – один канал занят (обслуживает заявку), очереди нет, S_2 – два канала занято, ..., S_k – занято k каналов, остальные свободны; ...; S_n – заняты все n каналов (очереди нет); S_{n+1} – заняты все каналы и 1 заявка в очереди, ...; S_{n+m} – заняты все n каналов, m заявок в очереди.

Показатели эффективности работы многоканальной СМО с ограниченной очередью:

1) Предельные вероятности:

$$P_0 = \left(1 + \frac{\rho}{1!} + \dots + \frac{\rho^n}{n!} + \frac{\rho^{n+1} (1 - (\rho/n)^m)}{n \cdot n! (1 - \rho/n)} \right)^{-1}, \dots, \quad (5.52)$$

$$P_k = \frac{\rho^k}{k!} P_0, \dots, P_n = \frac{\rho^n}{n!} P_0, \dots, P_{n+r} = \frac{\rho^{n+r}}{n^r \cdot n!} P_0$$

2) Вероятность отказа:

$$P_{отк} = P_{n+m} = \frac{\rho^{n+m}}{n^m \cdot n!} P_0 \quad (5.53)$$

3) Относительная пропускная способность:

$$Q = 1 - P_{отк} = 1 - \frac{\rho^{n+m}}{n^m \cdot n!} P_0 \quad (5.54)$$

4) Абсолютная пропускная способность:

$$A = \lambda Q = \lambda \left(1 - \frac{\rho^{n+m}}{n^m \cdot n!} P_0 \right) \quad (5.55)$$

5) Среднее число заявок в очереди:

$$L_{оч} = \frac{\rho^{n+1} \left[1 - (\rho/n)^m (m+1 - m\rho/n) \right]}{n \cdot n! (1 - \rho/n)^2} P_0 \quad (5.56)$$

6) Среднее число занятых каналов:

$$\bar{k} = \rho \left(1 - \frac{\rho^{n+m}}{n^m \cdot n!} P_0 \right) \quad (5.57)$$

7) Среднее число заявок в системе:

$$L_{сист} = \bar{k} + L_{оч} \quad (5.58)$$

8) Среднее время пребывания заявки в системе или в очереди определяется по формулам Литтла:

$$T_{сист} = \frac{L_{сист}}{\lambda}, \quad T_{оч} = \frac{L_{оч}}{\lambda} \quad (5.59)$$

ГЛАВА 6. ОСНОВЫ ТЕОРИИ ИГР

6.1. Основные понятия теории игр

При решении задач по теории оптимизации рассматривались такие задачи принятия решений, когда выбор решения осуществляется одним лицом. В подобных задачах рационального ведения хозяйства решение выбирается при предположении о том, что известны целевая функция, различные способы действия и ограничения. В дальнейшем рассмотрим задачи принятия решений в ситуациях с несколькими участниками, когда значение целевой функции для каждого из субъектов зависит от решений, принимаемых всеми остальными участниками.

Данные задачи рассматриваются в рамках теории игр. Предметом теории игр являются такие ситуации, в которых важную роль играют конфликты и совместные действия.

Конфликт может возникнуть также из различия целей, которые отражают не только не совпадающие интересы различных сторон, но и многосторонние интересы одного и того же лица.

Всякая претендующая на адекватность математическая модель социально-экономического явления должна отражать присущие ему черты конфликта, т.е. описывать:

а) множество заинтересованных сторон (обычно их называют *игроками*, они еще именуются субъектами, лицами, сторонами, участниками). В случае если число игроков конечно, они различаются по своим номерам (1-й игрок, 2-й игрок и т.д.), или по присваиваемым им именам (например, Продавец и Покупатель в ситуации монополии – монопсонии);

б) возможные действия каждый из сторон, именуемые также *стратегиями* или *ходами*;

в) интересы сторон, представленные функциями выигрыша (платежа) для каждого из игроков.

Можно сказать, что теория игр лежит на стыке общественных и математических наук.

Определение 6.1. Теория игр – это наука, изучающая решения людей, фирм, правительств и других агентов.

Теория игр — это раздел математики, в котором исследуются математические модели принятия решений в условиях конфликта, т. е. в условиях столкновения сторон, каждая из которых стремится воздействовать на развитие конфликта в своих собственных интересах.

Определение 6.2. Исход конфликта называется выигрышем.

В теории игр предполагается, что функции выигрыша и множество стратегий, доступных каждому из игроков, общеизвестны, т.е. каждый игрок знает свою функцию выигрыша и набор имеющихся в его распоряжении стратегий, а также функции выигрыша и стратегии всех остальных игроков, и в соответствии с этой информацией организует свое поведение.

Определение 6.3. Формализация содержательного описания конфликта представляет собой его математическую модель, которую называют *игрой*.

Выбор и осуществление одного из предусмотренных правилами действий называется ходом игрока. Ходы могут быть личными и случайными.

Определение 6.4. Личный ход – это сознательный выбор игроком одного из возможных действий.

Определение 6.5. Случайный ход – это случайно выбранное действие (например, выбор карты из перетасованной колоды)

В дальнейшем мы будем рассматривать только личные ходы игроков.

Определение 6.6. Игры, в которых есть личные ходы, называются стратегическими.

Определение 6.7. Стратегией игрока называется совокупность правил, определяющих выбор его действия при каждом личном ходе в зависимости от сложившейся обстановки.

Обычно в процессе игры при каждом личном ходе игрок делает выбор в зависимости от конкретной ситуации. Однако, в принципе, возможно, что все решения приняты игроком заранее. Это означает, что игрок выбрал определенную стратегию, которая может быть задана в виде списка правил или программы.

Самая известная задача из теории игр – это дилемма заключенного. Данная задача была сформулирована Мерилом Фладом и Мелвином Дresherом в 1950 году. Название дилемме дал математик Альберт Такер.

Классическая формулировка дилеммы заключённого такова:

Двое преступников — Джон и Билл — были пойманы полицией за совершенное преступление. Есть основания полагать, что они действовали по сговору, и полиция, изолировав их друг от друга, предлагает им одну и ту же сделку: если один свидетельствует против другого, а тот хранит молчание, то первый освобождается за помощь следствию, а второй получает максимальный срок лишения свободы (10 лет). Если оба молчат, их деяние проходит по более лёгкой статье, и каждый из них приговаривается к полугоду тюрьмы. Если оба свидетельствуют друг против друга, они получают минимальный срок (по 2 года). Каждый заключённый выбирает, молчать или свидетельствовать против другого. Однако ни один из них не знает точно, что сделает другой. Что произойдёт?

Игру можно представить с помощью так называемой платежной матрицы (табл. 6.1). Первое число в паре чисел – срок, который получит Джон, второе число – срок, который получит Билл.

Таблица 6.1.

		Стратегии Билла	
		Молчать	Сотрудничать
Стратегии Джона	Молчать	(2; 2)	(10; 0)
	Сотрудничать	(0; 10)	(5; 5)

Дилемма появляется, если предположить, что оба заботятся только о минимизации собственного срока заключения. Представим рассуждения одного из узников. Если партнёр молчит, то лучше его предать и выйти на свободу (иначе — два года тюрьмы). Если партнёр свидетельствует, то лучше тоже свидетельствовать против него, чтобы получить 5 лет (иначе — 10 лет) тюрьмы. Стратегия «свидетельствовать» строго доминирует над стратегией «молчать». Аналогично другой узник приходит к тому же выводу.

С точки зрения группы (этих двух узников) лучше всего сотрудничать друг с другом, хранить молчание и получить по два года, так как это уменьшит суммарный срок заключения.

Другой пример классической задачи из теории игр – трагедия общин. Термин «tragedy of the commons» появился из притчи Уильяма Форстера Ллойдагуен в его книге 1833 года о населении. Затем термин популяризировал Гарретт Хардингуен в 1968 году в статье для журнала Science.

Трагедия общин – род явлений, связанных с противоречием между интересами индивидов относительно блага общего пользования. В основном под этим подразумевается проблема истощения такого блага. В общем случае «трагедия» состоит в том, что свободный доступ к экономическому ресурсу (например, пастбищу) уничтожает или истощает ресурс из-за чрезмерного его использования. Это происходит потому, что все пользующиеся им получают выгоды непосредственно, но издержки содержания ресурса по каким-либо причинам вменить какому-то члену общины невоз-

можно, и/или они в той или иной степени возлагаются на всех членов общины.

6.2. Классификация игр

Различные виды игр можно классифицировать, основываясь на том или ином принципе:

- по числу игроков;
- по числу стратегий;
- по свойствам функций выигрыша;
- по возможностям предварительных переговоров;
- взаимодействия между игроками в ходе игры.

По числу игроков различают игры с двумя, тремя и более участниками. В принципе возможны так же игры с бесконечным числом игроков.

По количеству стратегий различают конечные и бесконечные игры. В конечных играх игроки располагают конечным числом возможных стратегий. Сами стратегии в конечных играх нередко называют чистыми стратегиями.

В бесконечных играх игроки имеют бесконечное число возможных стратегий, так в ситуации Продавец-Покупатель каждый из игроков может назвать любую цену и количество продаваемого (покупаемого) товара.

По свойствам функции выигрыша (платежных функций) – возможна ситуация, когда выигрыш одного из игроков равен проигрышу другого, т.е. налицо конфликт между игроками. Подобные игры называются **играми с нулевой суммой**, или **антагонистическими играми**. Прямой противоположностью такого типа являются игры с *постоянной разностью*, в которых игроки и выигрывают, и проигрывают одновременно так, что им выгодно действовать сообща.

Между этими крайними случаями имеется множество игр с ненулевой суммой, где имеются и конфликты, и согласованные действия игроков.

В зависимости от возможности предварительных переговоров между игроками различают кооперативные и некооперативные игры.

Игра называется *кооперативной*, если до начала игры игроки образуют коалиции и принимают взаимобязывающие соглашения о своих стратегиях.

Игра, в которой игроки не могут координировать свои стратегии подобным образом, называется некооперативной. Очевидно, что все антагонистические игры могут служить примером некооперативных игр.

6.3. Формальное представление игр

Множество всех игроков, обозначаемое I , в случае их конечного числа может задаваться простым перечислением игроков. Например, $I=(1,2)$ при игре в орлянку; $I=(1,2,\dots,n)$ в случае анализа результатов голосования.

Множество стратегий игрока i обозначим через x_i . При игре в орлянку каждый игрок располагает двумя стратегиями: $x_i=(орел, решка)$. Каждый участник голосования имеет выбор на множестве стратегий [За, Против].

Покупатель и продавец могут назначать некоторую неотрицательную цену на продаваемый (покупаемый) товар, т.е. множество стратегий каждого из них: x_i ; $p_0 > 0$. В каждой партии игрок выбирает свою стратегию x_i , в результате чего оказывается набор стратегий $x=(x_1, x_2, \dots, x_n)$, называемых *ситуацией*.

Ситуация в парламенте описывает список (за, за, против, за,...), полученный в итоге голосования.

Заинтересованность игроков в ситуациях проявляется в том, что каждому игроку i в каждой ситуации x приписывается число, выражающее степень удовлетворения его интересов в данной ситуации. Это число назы-

вается выигрышем игрока i и обозначается через $hi(x)$, а соответствие между набором ситуаций и выигрышем I называется **функцией выигрыша** (платежной функцией) этого игрока Hi .

В случаях конечной игры двух лиц функции выигрыша каждого из игроков удобно представить в виде **матрицы выигрышей**, где строки представляют стратегии одного игрока, столбцы – стратегии другого игрока, а в клетках матрицы указывают выигрыши каждого из игроков в каждой из образующихся ситуаций. (Данная форма представления конечных игр двух лиц объясняет общее для них название – **матричные игры**).

Например, в случае игры в орлянку каждый из игроков имеет по две стратегии Орел и Решка. Если оба игрока выбирают одинаковые стратегии, первый игрок выигрывает 1 рубль, а второй проигрывает рубль. В случае, когда оба игрока выбирают различные стратегии, первый игрок проигрывает один рубль, второй выигрывает один рубль.

Матрица выигрышей первого игрока

Стратегии первого иг- рока		Стратегии второго игрока		
		Орел	Решка	H_1
	Орел	1	-1	
	Решка	-1	1	

Матрица выигрышей второго игрока

Стратегии первого иг- рока		Стратегии второго игрока		
		Орел	Решка	H_2
	Орел	-1	1	
	Решка	1	-1	

Для антагонистических игр выполняется соотношение $H_1 = -H_2$

Для наглядности обе матрицы совмещают в одну:

Стратегии первого иг- рока		Стратегии второго игрока	
		Орел	Решка
	Орел	$\begin{pmatrix} (1;-1) & (-1;1) \\ (-1;1) & (1;-1) \end{pmatrix}$	
	Решка		

Рассмотрим пример записи функции для бесконечной игры. Для случая двух игроков, каждый из игроков может объявить цену p_i , по которой он хотел бы продать некоторое количество товара. При этом предполагается, что потребители приобретают товар у фирмы, объявившей меньшую цену, или распределят свой спрос между фирмами в случае, если они назначили одинаковую цену.

Если функцию спроса в зависимости от цены на товар обозначить – $d(p)$, то функция выигрыша первой фирмы $\Pi_1(p_1, p_2)$ будет выглядеть:

$$\Pi_1(p_1, p_2) = \begin{cases} p_1 d(p_1), & \text{если } p_1 < p_2 \\ p_1 \frac{d(p_1)}{2}, & \text{если } p_1 = p_2 \\ 0, & \text{если } p_1 > p_2 \end{cases}$$

6.4. Принцип решения матричных антагонистических игр

В качестве основного допущения в теории игр предполагается, что каждый игрок стремится обеспечить себе максимально возможный выигрыш при любых действиях партнера. Предположим, что имеется конечная антагонистическая игра с матрицей выигрышей первого игрока H_1 и соответственно матрицей выигрышей второго игрока H_2 . Пусть Игрок 1 считает, что, какую бы стратегию он ни выбрал, Игрок 2 выберет стратегию, максимизирующую его выигрыш, и тем самым минимизирующую выигрыш Игрока 1.

Таким образом, Игрок 1 выбирает i -ю стратегию, которая является решением задачи:

$$\max_i \min_j h_{ij}$$

Игрок 2 точно также стремится обеспечить себе наивысшую величину выигрыша (или, что эквивалентно, наименьшую величину проигрыша) вне зависимости от выбранной стратегии противника. Его оптимальной стратегией будет столбец H_0 с наименьшим максимальным платежом. Таким образом, Игрок 2 выберет j -ю стратегию, которая является решением задачи:

$$\min_j \max_i h_{ij}$$

В итоге, если Игрок 1 придерживается избранной стратегией (называемой **максиминной стратегией**), его выигрыш в любом случае будет не меньше максиминного значения (называемого «**нижней ценой игры**»), т.е.

$$h_{ij} \geq \max_i \min_j h_{ij}$$

Соответственно, если Игрок 2 придерживается своей минимаксной стратегии, то его проигрыш будет не больше максиминного значения (называемого «**верхней ценой игры**»), т.е.

$$h_{ij} \leq \min_j \max_i h_{ij}$$

В случае, когда верхняя цена игры равна нижней, т.е. $\max_i \min_j h_{ij} = \min_j \max_i h_{ij}^*$, оба игрока получают свои гарантированные платежи, а значение h_{ij}^* называется **ценой игры**.

Элемент матрицы h_{ij} матрицы выигрышей, соответствующей стратегиям, называется **седловой точкой матрицы H** .

В случае, если цена антагонистической игры равна 0, игра называется **справедливой**.

Рассмотрим игру, в которой Игрок 1 располагает двумя стратегиями, а Игрок 2 – тремя. Матрица выигрышей Игрока 1 имеет вид:

$$H_1 = \begin{pmatrix} 2 & 1 & 4 \\ -1 & 0 & 6 \end{pmatrix}$$

Поскольку мы рассматриваем пример антагонистической игры, то матрица выигрышей Игрока 2 будет $H_2 = -H_1$.

Игрок 1 рассчитывает, что если он выберет первую стратегию (т.е. первую строку матрицы H_1), то противник выберет свою вторую стратегию (т.е. второй столбец) так, что выигрыш будет равен 1. Если же он выбирает вторую стратегию, то противник может выбрать первую стратегию, так что выигрыш будет равен -1.

Проанализировав полученные значения: Игрок 1 останавливается на своей первой стратегии, которая обеспечивает ему максимальный гарантированный выигрыш, равный 1.

Точно также Игрок 2 рассматривает свои наихудшие варианты, когда противник выбирает первую или вторую стратегии, или когда противник выбирает вторую стратегию, когда Игроком 2 выбран третий столбец. Этим варианты соответствуют максимальным значениям столбцов 2, 1 и 6.

Взяв минимальные значения этих максимумов, Игрок 2 останавливается на своей второй стратегии, при которой его проигрыш минимален и равен 1:

$$\begin{aligned} & \min_j h_{ij} \\ & \begin{pmatrix} 2 & 1 & 4 \\ -1 & 0 & 6 \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} = \max_i \min_j h_{ij} \\ & \max_i h_{ij} \begin{pmatrix} 2 & 1 & 4 \\ -1 & 0 & 6 \end{pmatrix} = \min_j \max_i h_{ij} \end{aligned}$$

Следовательно, в этой игре существуют совместные выборы стратегий, т.е.

$$h^* = \max_i \min_j h_{ij} = \min_j \max_i h_{ij} = 1.$$

Следовательно, в этой игре разумно ожидать, что противники будут придерживаться избранных стратегий. Матричная антагонистическая игра, для которой $\max_i \min_j h_{ij} = \min_j \max_i h_{ij}$ - называется вполне определенной, или игрой, имеющей решение в чистых стратегиях.

Однако не все матричные антагонистические игры являются вполне определенными.

Игры, в которых выполняется строгое неравенство, называется не полностью определенными играми (или играми, не имеющими решения в чистых стратегиях).

Рассмотрим пример такой игры:

$$\begin{pmatrix} 6 & -2 & 3 \\ -4 & 5 & 4 \end{pmatrix}$$

Для этой игры $\max_i \min_j h_{ij} = -2 < 4 = \min_j \max_i h_{ij}$.

В итоге если игроки будут следовать предложенным выше правилам, то Игрок 1 выберет стратегию 1 и будет ожидать, что Игрок 2 выберет стратегию 2, при которой проигрыш равен -2, в то время как Игрок 2 выберет стратегию 3 и будет ожидать что Игрок 1 выберет стратегию 2 с выигрышем равным 4.

Однако если Игрок 2 выберет свою третью стратегию, то Игрок 1 поступит правильнее, выбирая вторую стратегию, а не первую стратегию. Аналогично, если Игрок 1 выберет первую стратегию, Игроку 2 выгоднее выбрать вторую стратегию, а не третью. По всей видимости, в играх подлобного типа принцип решения в чистых стратегиях оказывается непригодным.

В описанной ситуации игрокам становится важно, чтобы противник не угадал, какую стратегию он будет использовать. Для осуществления этого плана игрокам следует пользоваться так называемой смешанной стратегией.

По существу, смешанная стратегия игрока представляет собой схему случайного выбора чистой стратегии. Математически ее можно представить как вероятностное распределение на множестве чистых стратегий данного игрока. В итоге вектор $x = (x_1, x_2, \dots, x_m)$, где x_j соответствует вероятности применения Игроком 1 i -той стратегии и $\sum_i x_i = 1$, задает смешанную стратегию этого игрока. Аналогично определяется смешанная стратегия у Игрока 2 $y = (y_1, y_2, \dots, y_n)$.

Мы будем предполагать использование игроками их смешанных стратегий независимым, так что вероятность, с которой Игрок 1 выбирает i тую стратегию, а Игрок 2 - j -ю, равна $x_i y_j$. В этом случае платеж h_{ij} . Суммируя по i и j , найдем математическое ожидание выигрыша Игрока 1:

$$v_1 = \sum_i \sum_j x_i h_{ij} y_j$$

или матричных обозначениях

$$v_1 = xHy'.$$

На множестве смешанных стратегий Игрок 1, стремящийся достичь наибольшего из гарантированных выигрышей, выбирает вектор вероятностей X так, чтобы получить максимум минимальных значений ожидаемых выигрышей, т.е. он решает задачу:

$$v_1 = \max_x \min_y xHy'.$$

Аналогично целью Игрока 2 является достижение минимума максимальных значений своих проигрышей, т.е. он решает задачу

$$v_2 = \min_y \max_x x'Hy.$$

Фундаментальным результатом теории игр является так называемая Теорема о минимаксе, которая утверждает, что сформулированные задачи Игрока 1 и Игрока 2 всегда имеют решение для любой матрицы выигрышей H , и кроме того, $v_1 = v_2 = v$.

Как и для вполне определенных игр, стратегия x' Игрока 1 называется **Максиминной стратегией**, стратегия Игрока 2 y' - **минимаксной стратегией**, значение v - ценой игры; в случае, когда $v = 0$ игра называется справедливой.

Очевидным следствием из Теоремы о минимаксе является соотношение:

$$xHy' < \bar{v} < x'Hy.$$

которое означает, что никакая стратегия Игрока 1 не позволит выиграть ему сумму большую, чем цена игры, если Игрок 2 применит свою минимаксную стратегию, и никакая стратегия Игрока 2 не даст возможности проиграть ему сумму меньшую, чем цена игры, если Игрок 1 применяет свою максиминную стратегию.

Это верно также и для чистых стратегий, как для частного случая смешанных стратегий (т.к. чистая стратегия – это стратегия, используемая с вероятностью 1). Использование любой чистой стратегии, в случае если противник использует свою оптимальную стратегию, не позволяет выиграть больше (проиграть меньше) цены игры.

Это факт часто используют для разработки конкретных алгоритмов решения антагонистических матричных игр.

6.5. Решение матричных антагонистических игр

Вычисление оптимальных стратегий значительно усложняется с ростом числа стратегий. Для поиска оптимальных стратегий можно использовать несколько подходов.

Для уменьшения размерности игры используется доминирование строк и столбцов. Обычно говорят, что k -я строка матрицы H доминирует

i -ю строку (т. е. одна чистая строка доминирует другую), если $h_{ij} \leq h_{kj}$ для всех j , $h_{ij} < h_{kj}$ хотя бы для одного j .

Аналогично l -й столбец доминирует j -й столбец, если $h_{il} \geq h_{ij}$ для всех i , $h_{il} < h_{ij}$ хотя бы для одного i .

Смысл этого определения состоит в том, что доминирующая стратегия никогда не хуже, а в некоторых случаях даже лучше, чем доминируемая стратегия. Отсюда важный вывод – игроку нет необходимости использовать доминируемую стратегию. Это позволяет на практике все доминируемые строки и столбцы отбросить, что позволит уменьшить размеры матрицы (заметим, что этот подход может использоваться также при поиске решения в чистых стратегиях).

Пример. Рассмотрим игру со следующей матрицей:

$$\begin{pmatrix} 5 & 2 & 1 \\ 2 & 1 & 3 \\ 3 & 6 & 4 \end{pmatrix} \rightarrow \text{третья строка этой матрицы доминирует вторую}$$

Исключение второй строки приводит к матрице: $\begin{pmatrix} 5 & 2 & 1 \\ 3 & 6 & 4 \end{pmatrix} \rightarrow$ третий столбец в этой урезанной матрице доминирует второй, и исключение второго столбца дает: $\begin{pmatrix} 5 & 1 \\ 3 & 4 \end{pmatrix}$.

В итоге, если можно найти решение для полученной игры, то его легко использовать для решения исходной игры, просто прописав исключенным строкам и столбцам нулевые вероятности.

Другой метод упрощения матрицы основан на свойстве, согласно которому аффинное преобразование матрицы платежей (т.е. преобразование всех элементов матрицы H по правилу $h_{ij} = ah_{ij} + b$, где $a \neq 0$) не изменяет решения игры; кроме того, цена преобразованной игры v' может быть получена из цены первоначальной игры по такому же правилу: $v' = av + b$.

Это означает, что для задания игры в принципе безразлично, в каких единицах измеряются выигрыши (в рублях или долларах) прибавление (вычитание) некоторой фиксированной суммы b изменит на такую же сумму выигрыш (проигрыш) каждого из игроков не меняя решение игры.

Это свойство может быть использовано для упрощения и придания наглядности матрице выигрышей (использовано по аналогии с операциями над матрицами – умножение матрицы на постоянное число, сложение и вычитание строк, кроме того, это свойство позволяет любую матричную антагонистическую игру сделать справедливой, для этого необходимо вычислить цену игры из всех элементов матрицы выигрышей).

Кроме того может быть использован графический способ для решения игры 2×2 (и вообще игр $2 \times n$ или $m \times 2$).

Например, матрица выигрышей имеет вид:
$$\begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix}.$$

Пусть Игрок 1 выбирает свою первую стратегию с вероятностью x , а вторую с вероятностью $(1-x)$. Если Игрок 2 выбирает свою первую стратегию, то (из первого столбца матрицы) математическое ожидание для Игрока 1 будет равно $v_1 = x - (1-x) = 2x - 1$. Если Игрок 2 выбирает свою вторую стратегию, то в соответствии со вторым столбцом матрицы: $v_1 = -x + (1-x) = 1 - 2x$.

Каждое из этих уравнений может быть изображено графически отрезком прямой линии в области $[0,1]$ на графике с координатами x и v_1 .

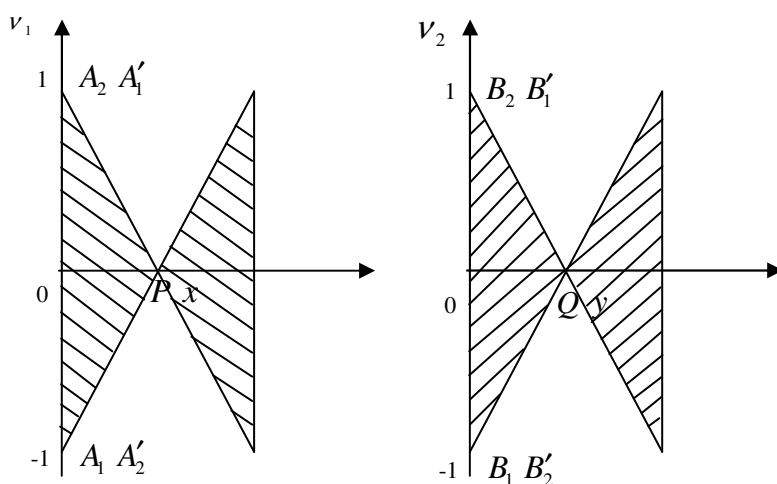


Рис. 6.1

Они представлены в виде линий $A_1A'_1$ и $A_2A'_2$, которые пересекаются в точке P . При данном x Игрок 1 может получить две величины v_1 , если Игрок 2 применяет свои чистые стратегии. Промежуточные значения v_1 соответствующие точкам между этими графиками, получаются если Игрок 2 применяет смешенные стратегии. Меньшая из этих величин, соответствующая каждому значению x показывает тот минимум, который может получить Игрок 1, выбирая стратегию $(x, (x-1))$.

Следовательно, линия $A_1PA'_2$ показывает платеж, который Игрок 1 может гарантированно получить при любой стратегии Игрока 2. Игрок 1 выбирает такое значение x , чтобы достичь наивысшей точки. Как показано на графике этой точкой является точка P , для которой $x = 0,5$; $v_1 = 0$.

Аналогично можно проанализировать игру Игрока 2, который использует свою первую стратегию с вероятностью y , а вторую с вероятностью $(1-y)$. На графике рядом ломанная $B_2QB'_1$ представляет наибольший проигрыш Игрока 2, при различных выборах y . Игрок 2 выбирает y так, чтобы достичь низшей точки на этой линии, т.е. точки Q , для которой $y = 0,5$; $v_2 = 0$. Следовательно, игроки должны применять свои стратегии с одинаковой вероятностью равной 0,5 и цена игры при этом будет $v_1 = v_2 = 0$, т.е. описанная игра справедлива.

Это позволяет сформулировать следующую задачу для Игрока 1:

$$\left\{ \begin{array}{l} h_{11}x_1 + h_{12}x_2 + ... + h_{1m}x_n \leq v \\ h_{21}x_1 + h_{22}x_2 + ... + h_{2n}x_n \leq v \\ \\ h_{m1}x_1 + h_{m2}x_2 + ... + h_{mn}x_n \leq v \\ x_1 + x_2 + ... + x_n = 1, \quad x_i \geq 0, \quad (i = \overline{1,n}) \end{array} \right.$$

Таким образом, в общем случае для решения матричной антагонистической игры размерностью $m \times n$ необходимо решить пару двойственных задач линейного программирования, в результате чего найдется набор оптимальных стратегий x^*, y^* и цены v .

ГЛАВА 7. ПРИНЯТИЕ РЕШЕНИЙ В УСЛОВИЯХ НЕОПРЕДЕЛЕННОСТИ И РИСКА

7.1. Математические методы оценки риска

В теории принятия решений выделяются два типа моделей: 1) **принятие решений в условиях риска**, когда лицо, принимающее решение, знает вероятности наступления исходов или последствий для каждого решения; 2) **принятие решения в условиях неопределенности**, когда лицо, принимающее решение, не знает вероятности наступления исходов или последствий для каждого решения. Исходная информация для принятия решения в ситуациях как неопределенности, так и риска обычно представляется с помощью таблицы выплат.

Под ситуацией риска, как уже отмечалось, в теории принятия решений понимается такая ситуация, когда можно указать не только возможные последствия каждого варианта принимаемого решения, но и вероятности их появления. Для выбора оптимального решения в данном случае предназначены критерий математического ожидания и критерий Лапласа.

Критерий математического ожидания является основным для принятия решения в ситуации риска. Он рассчитывается следующим образом:

$$K = \max_{i=1, \dots, m} \{M_i\} \quad (7.1)$$

$$M_i = \sum_{j=1}^n x_{ij} \cdot p_j \quad (7.2)$$

где x_{ij} – выплата, которую можно получить в j -м состоянии «среды»; p_j – вероятность j -го состояния среды. Используя данный критерий, лучшей стратегией будет та, которая обеспечивает наибольший средний выигрыш.

Рассмотри следующую задачу. Пусть первым игроком является банк, а вторым игроком или «природой» – ситуация на валютном рынке. Страте-

гии банка – это возможности покупки различных валют: А1 – банк покупает доллар США; А2 – банк покупает евро; А3 – банк покупает английский фунт стерлингов. Возможные ситуации на валютном рынке (стратегии «природы»): П1 – курс валюты по отношению к рублю упадет; П2 – курс валюты по отношению к рублю не изменится; П3 – курс валюты по отношению к рублю возрастет. Прибыль, которую получит банк, благодаря курсовой разнице покупки и продажи валюты, из расчета на 1 вложенный рубль, задана платежной матрицей (см. таблица 7.1). Задача заключается в определении такой стратегии (чистой или смешанной), которая при ее применении обеспечила бы оперирующей стороне наибольший выигрыш.

Таблица 7.1.

Стратегии	Валютный рынок		
	В1 ($p_1 = 0,5$)	В2 ($p_1 = 0,2$)	В3 ($p_1 = 0,3$)
А1	0,2	0,8	0,7
А2	1,5	3,0	0,6
А3	1,1	2,5	2,0

Рассчитаем математические ожидания для каждой из стратегий:

$$M_1 = 0,2 \cdot 0,5 + 0,8 \cdot 0,2 + 0,7 \cdot 0,3 = 0,47$$

$$M_2 = 1,5 \cdot 0,5 + 3,0 \cdot 0,2 + 0,6 \cdot 0,3 = 1,53$$

$$M_3 = 1,1 \cdot 0,5 + 2,5 \cdot 0,2 + 2,0 \cdot 0,3 = 1,65$$

Таким образом, наилучшей стратегией будет покупка английского фунта, для которой ожидается наибольшая прибыль.

Критерий Лапласа. Если ни один из возможных вариантов нельзя назвать наиболее вероятным по сравнению с остальными, то решение можно принимать с помощью критерия Лапласа:

$$L = \max_{i=1,\dots,m} \left\{ \sum_{j=1}^n x_{ij} \right\} \quad (7.3)$$

На основании приведенной формулы оптимальным надо считать то решение, которому соответствует наибольшая сумма выплат. Сумма выплат для отдельных вариантов решений в нашем примере составит:

$$\sum_{j=1}^3 x_{1j} = 1,7$$

$$\sum_{j=1}^3 x_{2j} = 5,1$$

$$\sum_{j=1}^3 x_{3j} = 5,6$$

Наибольшей является сумма выплат третьей строки. Значит, в качестве оптимального решения надо покупку английских фунтов, т.е. то же решение, что было признано оптимальным и с помощью критерия математического ожидания.

Когда два разных критерия предписывают принять одно и то же решение, то это лишний раз подтверждает его оптимальность. Если же критерии указывают на разные решения, то предпочтение в ситуации риска надо отдать тому из них, на которое указывает критерий математического ожидания. Именно он является основным для данной ситуации, однако существуют и другие.

Часто на практике в силу ряда причин оценить вероятность того или иного исхода невозможно. Такая ситуация называется ситуацией неопределенности. Для выбора эффективной стратегии в ситуации неопределенности используются следующие критерии.

Критерий MAXIMAX. Он определяет альтернативу, максимизирующую максимальный результат для каждого состояния возможной действительности. Это критерий крайнего оптимизма. Наилучшим признается решение, при котором достигается максимальный выигрыш, равный

$$M = \max_{i=1,\dots,m} \left(\max_{j=1,\dots,n} x_{ij} \right) \quad (7.4)$$

Из данной формулы видно, что сначала осуществляется поиск максимального значения по столбцам, а затем среди найденных значений также находят максимум. В нашем примере наилучшим вариантом будет вторая стратегия банка – покупка евро.

Однако применение такого критерия на практике является довольно рискованным.

Максиминный критерий Вальда. Другое название данного критерия критерий пессимиста, поскольку при его использовании как бы предполагается, что от любого решения надо ожидать самых худших последствий.

$$W = \max_{i=1,\dots,m} \left(\min_{j=1,\dots,n} x_{ij} \right) \quad (7.5)$$

Расчет максимина в соответствии с приведенной выше формулой состоит из двух шагов. Находим худший результат каждого варианта решения, т.е. величину $\min_{j=1,\dots,n} x_{ij}$ и строим таблицу (табл. 7.2).

Таблица 7.2.

Стратегии	Столбец минимумов
A1	0,2
A2	0,6
A3	1,1

Из худших результатов, представленных в столбце минимумом, выбираем лучший. Он относится к третьей строке таблицы, что предписывает покупку английских фунтов.

Это перестраховочная позиция крайнего пессимиста. Такая стратегия приемлема, когда инвестор не столько заинтересован в крупной удаче, сколько хочет застраховать себя от неожиданных проигрышей. Выбор такой стратегии определяется отношением принимающего решения лица к риску.

Критерии MINIMAX, или критерий Сэвиджа. И отличие от предыдущего критерия он ориентирован не столько на минимизацию потерь, сколько на минимизацию сожалений по поводу упущенной прибыли. Он допускает разумный риск ради получения дополнительной прибыли. Пользоваться этим критерием для выбора стратегии поведения в ситуации неопределенности можно лишь тогда, когда есть уверенность в том, что случайный убыток не приведет фирму (проект) к полному краху.

$$S = \min_{i=1,\dots,m} \left\{ \max_{j=1,\dots,n} \left\{ \max_{i=1,\dots,m} x_{ij} - x_{ij} \right\} \right\} \quad (7.6)$$

Расчет данного критерия включает в себя четыре шага.

1. Находим лучшие результаты каждого в отдельности столбца, т.е. $\max_{i=1,\dots,m} x_{ij}$. Таковыми в нашем примере будут: для первого столбца – 1,5, для второго – 3,0 и третьего – 2,0. Это те максимумы, которые можно было бы получить, если бы удалось точно угадать возможные реакции рынка.

2. Определяем отклонения от лучших результатов в пределах каждого отдельного столбца, т.е. $\max_{i=1,\dots,m} x_{ij} - x_{ij}$. Получаем матрицу отклонений, которую можно назвать матрицей сожалений, поскольку её элементы – это недополученная прибыль от неудачно принятых решений из-за ошибочной оценки возможной реакции рынка. Матрицу сожалений можно оформить в виде таблицы (табл. 7.3).

Таблица 7.3. Матрица сожалений

Стратегии	Возможные размеры упущенной прибыли		
	B1	B2	B3
A1	1,3	2,2	1,3
A2	0	0	1,4
A3	0,4	0,5	0

Таблица 7.4. Максимальные сожаления

Стратегии	Столбец минимумов
A1	2,2
A2	1,4
A3	0,5

3. Для каждого варианта решения, т.е. для каждой строки матрицы сожаления, находим наибольшую величину. Получаем столбец максимумов сожалений (табл. 7.4).

4. Выбираем решение, при котором максимальное сожаление будет меньше других. В приведенном столбце максимальных сожалений оно стоит в третьей строке.

Критерий пессимизма-оптимизма Гурийца. При выборе решения он рекомендует руководствоваться некоторым средним результатом, характеризующим состояние между крайним пессимизмом и безудержным оптимизмом, т.е. критерий выбирает альтернативу с максимальным средним результатом (при этом действует негласное предположение, что каждое из возможных состояний среды может наступить с равной вероятностью). Формально данный критерий выглядит так:

$$H = \max_{i=1,...,m} \left\{ k \cdot \min_{j=1,...,n} x_{ij} + (1 - k) \cdot \max_{j=1,...,n} x_{ij} \right\}, \quad (7.7)$$

где k – коэффициент пессимизма, который принадлежит отрезку от 0 до 1 в зависимости от того, как принимающий решение оценивает ситуацию. Если он рассматривает данную ситуацию как оптимистичную, то коэффициент должен быть меньше 0,5. При $k = 0$ критерий Гурвица совпадает с максимальным критерием, а при $k = 1$ – с критерием Вальда.

Рассчитаем критерий Гурвица для нашего примера, приняв критерий пессимизма, равным 0,6.

$$H_1 = 0,6 \cdot 0,8 + 0,4 \cdot 0,2 = 0,56$$

$$H_2 = 0,6 \cdot 3,0 + 0,4 \cdot 0,6 = 2,04$$

$$H_3 = 0,6 \cdot 2,5 + 0,4 \cdot 1,1 = 1,95$$

Второе максимальное значение критерия Гурвица означает, что следует выбрать стратегию A2.

7.2. Оценка риска с помощью VaR

Valueatrisk (VaR) – это максимальная сумма денег, которую может потерять портфель инвестора в течение заданного промежутка времени с заданной вероятностью (доверительной вероятностью) на «нормальном» рынке.

«Нормальный» рынок – это динамика рыночных факторов в отсутствии системного кризиса в экономике или негативных фактов (событий), способных вызвать существенное изменение факторов.

VaR – это стоимость под риском. В основе расчета – динамика доходностей (изменения цен) инструмента на заданном временном интервале.

VaR состоит из трех компонентов.

1) Доверительная вероятность (уровень доверия). По базельским документам используется величина 99 %, в системе RiskMetrics – 95 %.

2) Временной горизонт, который зависит от рассматриваемой ситуации, по базельским документам — 10 дней, по методике RiskMetrics – 1

день. Чаще распространен расчет с временным горизонтом 1 день. 10 дней используется для расчета величины капитала, покрывающего возможные убытки.

3) Возможные потери (количество денег (обычно долларов) или процентах).

VaR отвечает на следующий вопрос: «Какой максимальный убыток можно ожидать в течение определенного отрезка времени с заданным уровнем вероятности»?

VaR позволяет оценить убытки с определенной вероятностью. И сделать это можно довольно кратко, чтобы человек мог относительно легко представить размер риска. Например, $VaR = 400$ тыс. руб. с уровнем доверия 99% может означать следующее:

с вероятностью 0,01 мы можем потерять 400 тыс. руб. и более в течение дня;

с вероятностью 99% мы не потеряем более 400 тыс. руб. в течение дня.

VaR отвечает на следующий вопрос: «Какой максимальный убыток я могу ожидать в течение определенного отрезка времени с заданным уровнем вероятности(доверия)»?

Чтобы рассчитать VaR, необходимо выяснить вид распределения доходностей портфеля инвестора. Обычно это нормальное распределение, логнормальное, Стюдента, Лапласа и н.др. На рисунке 7.1 схематично изображен геометрический смысл расчета VaR.

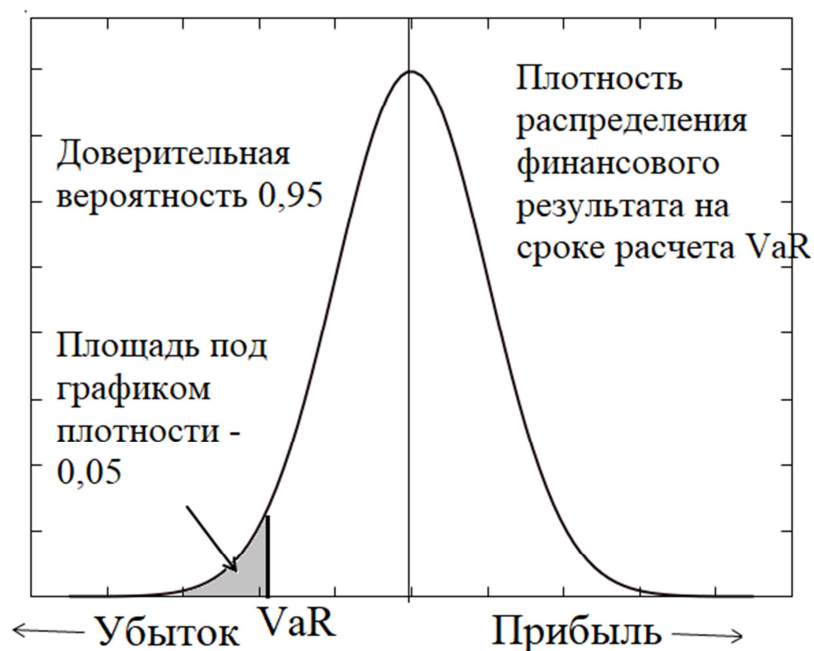


Рис. 7.1

В современной практике временной горизонт определяют обычно по одному из следующих двух критериев: намечаемого периода владения рассматриваемым активом (то есть времени его удержания в портфеле предприятия); уровня его ликвидности (реального срока его конверсии в денежную форму без потери своей текущей рыночной стоимости).

Преимуществом показателя VaR является то, что это довольно простая для понимания и интерпретации мера интегрального риска по портфелю, которая учитывает влияние сразу всех рыночных факторов.

Недостатки VaR:

- 1) неаддитивная мера;
- 2) использует исторические данные – учитывает прошлые события и динамику рынка;
- 3) не оценивает величину потерь в хвосте распределения доходностей;
- 4) обычно не учитывает ликвидность инструментов;
- 5) есть модельный риск.

Существует три группы методов получения VaR.

1. Непараметрические: исторические, бутстрэп, взвешенные исторические подходы, историческое моделирование, взвешенное по возрасту, историческое моделирование, взвешенное по волатильности, историческое моделирование, взвешенное по корреляции и фильтрованное историческое моделирование.

2. Параметрические - параметрический метод для изолированного актива, параметрический метод для многокомпонентного портфеля (вариационно-ковариационный).

3. Метод Монте-Карло. Параметрический метод для одного актива. Базируется на предположении о нормальном распределении доходности. нормальном распределении соответствующей экономической характеристики (цены актива, валютного курса и т. д.). Пусть μ и σ — параметры нормального распределения (среднее значение и стандартное отклонение), оценённые по прошлым наблюдениям. Значение VaR, полученное для нормального распределения для одного вида актива:

$$VaR = S \cdot (M(r) - z_{1-\alpha} \cdot \sigma(r))$$

где S – текущее значение капитала; α – заданный уровень доверительной вероятности; $1 - \alpha$ — вероятность потерь; $z_{1-\alpha}$ – нижний квантиль стандартного нормального распределения (площадь под кривой функции плотности распределения левее $z_{1-\alpha}$, так называемая обратная функция стандартного нормального распределения). VaR в приведённой формуле равен квантилю (quantile-VaR), или верхней границе одностороннего доверительного интервала в «левом хвосте» распределения. Когда речь идёт о временном горизонте в один день, μ часто принимают равным нулю. В этом случае получим максимальное значение VaR:

$$VaR = S \cdot z_{1-\alpha} \cdot \sigma(r)$$

Библиографический список

1. Адамчук, А. С. Математические методы и модели исследования операций (краткий курс) : учебное пособие / А. С. Адамчук, С. Р. Амиров, А. М. Кравцов. — Ставрополь : СКФУ, 2014. — 163 с. — Текст : электронный // Лань : электронно-библиотечная система. — URL: <https://e.lanbook.com/book/155283>.
2. Бережная Е.В., Бережной В.И. Математические методы моделирования экономических систем: Учеб. пособие. — 2-е изд., перераб. и доп. — М.: Финансы и статистика, 2006. — 432 с.
3. Бочаров, П.П. Теория массового обслуживания : Учебник. / П.П. Бочаров, А.В. Печинкин – М. : Изд-во РУДН, 1995. – 529 с.
4. Бурда, А. Г. Исследование операций в экономике : учебное пособие / А. Г. Бурда, Г. П. Бурда. — Санкт-Петербург : Лань, 2021. — ISBN 978-5-8114-3149-6. — Текст : электронный // Лань : электронно-библиотечная система. — URL: <https://e.lanbook.com/book/169285>.
5. Гребенникова, И. В. Методы оптимизации : учебное пособие / И. В. Гребенникова. — Екатеринбург : УрФУ, 2017. — 148 с.
6. Исследование операций в экономике : Учебн. Пособие для вузов / Н.Ш. Кремер, Б.А. Путко, И.М. Тришин, М.Н. Фридман; по ред. Проф. Н.Ш. Кремера. – М. – ЮНИТИ. 2000. – 407 с.
7. Калмыков, С. И. Исследование операций : учебное пособие / С. И. Калмыков, М. А. Первухин, А. А. Степанова. — Владивосток : ВГУЭС, 2019. — 152 с. — ISBN 978-5-9736-0555-1. — Текст : электронный // Лань : электронно-библиотечная система. — URL: <https://e.lanbook.com/book/161484> (дата обращения: 28.02.2021).
8. Косоруков О.А., Мищенко А.В. Исследование операций: Учебник / Косоруков О.А., Мищенко А.В. // под общ. ред. д.э.н., проф. Н.П. Тихомирова. – М. : Издательство «Экзамен», 2003. – 448 с.

9. Кремер Н.Ш., Путко Б.А., Тришина И.М. Математика для экономистов : от Арифметики до Эконометрики : учеб.-справоч. Пособие / под ред. Проф. Н.Ш. Кремера. – М. : Высшее образование, 2007. – 646 с.

10. Майника, Э. Алгоритмы оптимизации на сетях и графах : пер. с англ. / Э. Майника. – М. : Мир, 1981. – 323 с.

11. Маховикова, Г.А. Анализ и оценка в бизнесе : учебник для академического бакалавриата / Г.А. Маховникова, Т.Г. Касьяненко. – М. : Издательство Юрайт, 2016. - 464 с.

12. Писарук, Н.Н. Исследование операций / Н.Н. Писарук. – Минск : БГУ, 2015. – 304 с.

13. Прокопенко Н. Ю. Методы оптимизации [Текст]: учеб. пособие / Н. Ю. Прокопенко; Нижегород. гос. архитектур. - строит. ун-т. – Н. Новгород: ННГАСУ, 2018. – 118 с.

14. Таха, Хэмди А. Введение в исследование операций. 6-е издание. : Пер. с англ. – М. : Издательский дом «Вильямс», 2001. – 912 с.

15. Уколов, А.И. Оценка рисков : учебник / А.И. Уколов. – 2-е изд. стер. – Москва : Директ-Медиа, 2018. – 627 с. : ил., схем., табл. – Режим доступа: по подписке. – URL: <http://biblioclub.ru/index.php?page=book&id=445268>.

Электронное учебное издание

Алексей Викторович Алпатов

ИССЛЕДОВАНИЕ ОПЕРАЦИЙ: КОНСПЕКТ ЛЕКЦИЙ (ЧАСТЬ 2)

Учебное пособие

Электронное издание сетевого распространения

Редактор Матвеева Н.И.

Темплан 2021 г. Поз. № 7.

Подписано к использованию 07.09.2021. Формат 60x84 1/16.

Гарнитура Times. Усл. печ. л. 6,6.

Волгоградский государственный технический университет.
400005, г. Волгоград, пр. Ленина, 28, корп. 1.

ВПИ (филиал) ВолгГТУ.
404121, г. Волжский, ул. Энгельса, 42а.