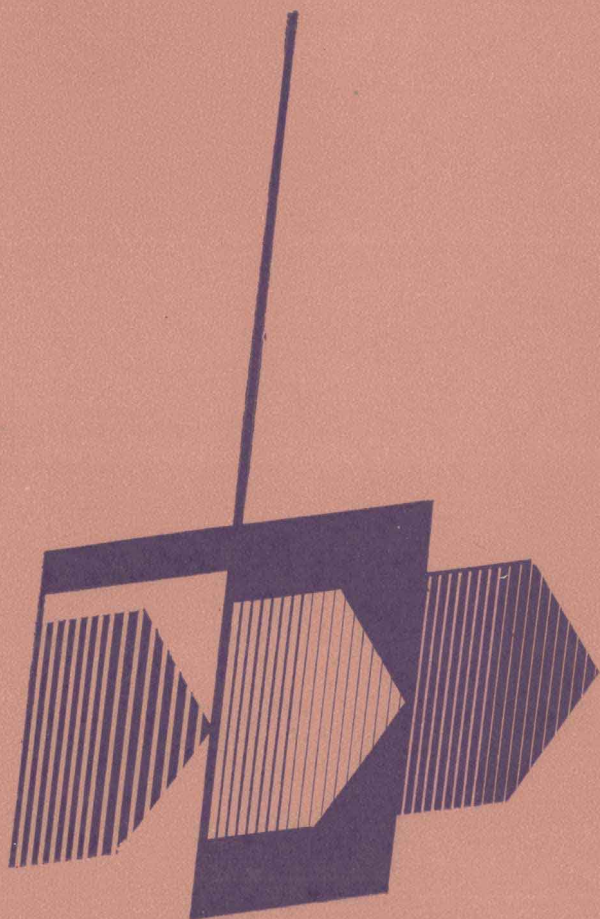


ТЕЛЕОБРАБОТКА ПЛАНОВО- ЭКОНОМИЧЕСКОЙ ИНФОРМАЦИИ



ГОСПЛАН МОЛДАВСКОЙ ССР
Научно-исследовательский институт планирования

**ТЕЛЕОБРАБОТКА
ПЛАНОВО—
ЭКОНОМИЧЕСКОЙ
ИНФОРМАЦИИ**

КИШИНЕВ „ШТИИНЦА“ 1985

В сборнике рассматриваются вопросы развития и использования фактографических информационных систем, алгоритм автоматического формирования запросов и метод синтеза программ расчета значений агрегативных показателей для задач АСУ. Приводятся модели прогнозирования, планирования и имитации развития народного хозяйства в интерактивном режиме, модели центрального вычислительного комплекса ВЦКП двухуровневой иерархической структуры.

Книга рассчитана на работников органов управления, специалистов, занимающихся разработкой и эксплуатацией автоматизированных систем, на аспирантов и студентов соответствующих специальностей.

Р е д к о л л е г и я

П.С.Солтан (ответственный редактор), С.В.Максимилиан (зам. ответственного редактора), М.А.Настас (ответственный секретарь), Е.Н.Гырла, С.К.Корня, В.П.Недельчук

С.К.Корня

ИНФОРМАЦИОННАЯ СИСТЕМА С МНОГОЯЗЫКОВЫМ СЕМАНТИЧЕСКИМ УРОВНЕМ

Тенденция построения баз данных с многоуровневой архитектурой, наметившаяся в последние годы в нашей стране и за рубежом, полностью оправдала себя. Сегодня имеется множество информационно-вычислительных систем с надстройкой в виде семантического уровня, использующих формализованные, близкие к естественному языку описания данных. Языки описания, а также языки общения (запросов) с базами данных, созданные на их основе, являются удобным средством для работы конечного пользователя. В большинстве случаев языки создаются путем формализации естественного русского языка и используются в различных информационных системах, функционирующих в СССР. Однако всестороннее развитие процессов экономической интеграции в рамках СЭВ, а также различных видов интернациональных взаимосвязей (торговля, научно-техническое сотрудничество и др.) ставят на повестку дня их автоматизацию. Предлагаемая работа представляет собой результат эксперимента по созданию средств автоматизации упомянутых процессов. Основой системы обработки информации подобного типа является проблемно-ориентированная база данных. Пользователями такой базы могут выступать представители различных государств. Следовательно, одним из обязательных требований к подобным системам должно быть обеспечение возможности идентичного семантического описания одних и тех же данных на языке того государства, представителем которого является пользователь. Очевидно, что создание таких средств сводится к многоязыковому семантическому описанию данных. Эксперимент проводился на информационной системе СПИРИТ, логическая схема организации которой дана в [1]. База данных организована разработанным нами программно-ассоциативным методом и представлена семантической, ассоциативной и собственной областями. При создании многоязыковых средств представления и описания данных ассоциативная и собственная области не претерпевают никаких изменений. Модификации подвергается лишь семантическая область, в частности словари наименований показателей и значений признаков (описательных терминов).

Словарь наименований показателей $S = \{s_l; l = \overline{1, l_0}\}$ представляется множеством $S = \{s_l^e; l = \overline{1, l_0}; e = \overline{1, e_0}\}$, где s - конкретное наименование отдельной позиции словаря; l - шифр наименования показателя; e - номер языка, на котором дается формализованная позиция наименования показателя или описательного термина.

Аналогично представляется словарь описательных терминов

$$Z = \{z_\beta; \beta = \overline{1, \beta_0}; e = \overline{1, e_0}\},$$

где z - конкретное наименование описательного термина отдельной позиции словаря; β - шифр описательного термина.

В общем случае $S = \bigcup z_l^e$; $Z = \bigcup z_\beta^e$.

Остальные моменты формального представления языка описания данных, равно как и логические процедуры вычислительного процесса обработки данных, достаточно полно описанные в [1; 2], остаются без изменения.

Эксперименты проводились на восьми языках (иностранных и народов СССР): русском, английском, французском, немецком, испанском, румынском, украинском и молдавском.

Процедуры обработки осуществляются по принципу многоаспектного избирательного поиска. Исходным является код запроса. Основные элементы кода следующие: номер класса показателей; позиции горизонтальной и вертикальной разверток; номера признаков; количество разворачиваемых позиций; начальная позиция развертки;

по аспектным признакам: номера признаков; количество разворачиваемых позиций; конкретные позиции на выбор;

по дополнительным признакам: номера признаков; конкретные позиции развертки;

по языковому признаку: общее количество языков, на которых должно выдаваться семантическое сопровождение результатов поиска; позиции конкретных языков при выборочном запросе.

Алгоритм поиска.

1. Начало.

2. Ввод кодового запроса.

3. Определение номера набора признаков искомого класса показателей.

4. Определение начальных позиций описательных терминов по всем признакам набора.

5. Селекция терминов требуемого языка описания выходных результатов.

6. Селекция строк ассоциативных матриц, соответствующих описательным терминам.

7. Определение адреса оснований показателей.

8. Селекция оснований по определенным адресам.

9. Формирование выходных результатов.

10. Конец.

В разработанной программе, реализующей приведенный укрупненный алгоритм, учтены лишь основные моменты многоаспектной селекции, наиболее приемлемые для справочного режима. Простота алгоритма и структуры организации базы данных позволяет реализовать множество версий, где могут быть учтены и другие, более детальные требования пользователей. Аналогичным образом разрабатываются алгоритмы обработки и формирования более сложных выходных структур.

Следует отметить некоторые затруднения при реализации системы. Почти каждый язык имеет свои алфавитные особенности, например, немецкий - $\ddot{u}, \ddot{a}$, румынский - $\hat{i}, \hat{s}, \hat{t}, \hat{a}$, молдавский - ж и др. Для данных символов не предусмотрены соответствующие коды, и аппаратура не может ни считывать их при вводе, ни печатать при выводе результатов на АШТУ или на экран видеотерминала. Представляется целесообразным более глубокое исследование особенностей алфавитов различных языков и учет их как при производстве аппаратурных вычислительных средств, так и при разработке программного обеспечения вычислительных систем.

Для иллюстрации получаемых результатов в терминах различных языков приведем описание единичного показателя на формализованных языках.

На русском языке:

ЗАТРАТЫ / СОЗДАНИЕ И ВНЕДРЕНИЕ АСУ И ВТ / МИНИСТЕРСТВО ЛЕГКОЙ ПРОМЫШЛЕННОСТИ / 1982 ГОД / ФАКТ / ТЫСЯЧ РУБЛЕЙ / - R_1 .

(R_1 - числовое основание искомого показателя).

Дробной чертой отделены друг от друга наименования показателя и описательные термины. В данном случае языком описания является язык с ограниченными грамматическими средствами, но пользователь легко преодолевает эти ограничения путем самостоятельного добавления недостающих союзов и необходимых изменений грамматических форм. Например, в описании показателя R_1 пользователь легко определит, что это затраты на создание АСУ и ВТ по министерству легкой промышленности за 1982 год, фактически в тысячах рублей.

На английском языке:

EXPENDITURES/CREATION AND INCULCATION OF ACS AND CALCULATION MACHINERY/MINISTRY OF LIGHT INDUSTRY/1982 YEAR/ FACT/THOUSAND ROUBLES/ - R_1 .

На французском языке:

LES DEPENSES/LA CREATION ET L'INCULCATION DU SAC ET DE LA TECHN.A
CALC./LE MINISTERE DE L'INDUSTRIE LEGERE/L'ANNEE 1982/LE FAIT/MIL-
LIES RUBLES/ - R₁.

На немецком языке:

AUSGABEN/SCHAFFUNG UND EINFUHRUNG DER ASS UND DER RECHENTECHNIK/MINISTERIUM DER LEICHTINDUSTRIE/1982 JAHR/FAKT/TAUSAND RUBEL/ - R₁.

На испанском языке:

GASTOS/CREACION E INTRODUCCION SAD Y TECNICA DE COMPUTACION/MINISTERIO DE LA INDUSTRIA LIVIANA/ANO 1982/FACTO/MIL RUBLOS/ - R₁.

На румынском языке:

CHeltuELI/CREAREA SI APLICAREA IN PRACTICA A SAC SI TEHNIC.DE CALCUL/MINISTERUL INDUSTRIEI USOARE/ANUL 1982/REAL/MII RUBLE/ - R₁.

На украинском языке:

ЗАТРАТИ/ СТВОРЕННЯ І ВПРОВАДЖЕННЯ АСУ І ВТ/ МІНІСТЕРСТВО ЛЕГКОЇ ПРОМИСЛОВОСТІ/ 1982 РІК/ФАКТ/ ТИС.КАРБ./ - R₁.

На молдавском языке:

КЕЛТУЕЛЬ/КРЕАРЯ ШИ АПЛИКАРЯ ЛН ПРАКТИКЭ А САК ШИ ТЕХНИЧИЙ ДЕ КАЛКУЛ/ МИНИСТЕРУЛ ИНДУСТРИЕЙ УШОАРЕ/ АНУЛ 1982/ РЕАЛ/ МИЙ РУБЛЕ/ - R₁.

Программная реализация позволяет осуществлять многоаспектное информационное обслуживание на любом языке или наборе языков на выбор как в пакетном режиме, так и в режиме теледоступа.

Л и т е р а т у р а

1. К о р н я С. К. Инструментарий построения безызыбочных баз данных. - В кн.: Информационные языки и базы данных в планировании. М., ЦУМИ АН СССР, 1980.

2. К о р н я С. К. Информационная совместимость взаимодействующих автоматизированных систем. - В кн.: Обработка информации в АСУ-Молдавия. Кишинев, Штиинца, 1980.

В.П.Недельчук, В.С.Гыля

АЛГОРИТМ СИНТЕЗА ЗАПРОСОВ
К СИСТЕМЕ ОБРАБОТКИ ДАННЫХ

В практике планирования и управления значительное число задач заключается в том, чтобы определить и в форме таблиц представить

численные значения показателей, характеризующих объекты управления, т.е. сводится или может быть сведено к заполнению заданных таблиц (форм) T числами, полученными в результате обработки другой информации или их поиска в различных фондах. Такое представление показателей привычно для широкого круга специалистов, и поэтому большинство задач, решаемых сегодня в рамках АСУ, является задачами автоматизации процесса получения именно такого рода таблиц.

Анализ структур различных форм представления показателей свидетельствует, что подлежащие заполнению таблицы T , как правило, состоят из следующих компонентов:

надтабличная текстовая часть A ;

табличная текстовая часть B ;

собственно числовая часть C .

Надтабличная текстовая часть A в общем случае содержит наименование таблицы (A_1), дополняющую информацию (A_2) и вспомогательную информацию (A_3). Содержание A_1 и/или A_2 иногда необходимо для семантического однозначного описания элементов-клеток (чисел) части C , причем A_1 задает наименование однородной совокупности показателей, а A_2 - конкретизирующие общие признаки данной совокупности (например, единицу измерения, рассматриваемый период и пр.); A_3 содержит номер таблицы, наименование органа, утвердившего ее, и тому подобную информацию, которая с точки зрения информативности алгоритма определения показателей T является лишней. Зачастую, однако, табличная текстовая часть B , заданная в виде вертикальных (B_1) и/или горизонтальных (B_2) граф, подграф (в дальнейшем - доменов), семантически однозначно описывает содержание всех элементов (чисел) C .

Таким образом, T можно представить в виде

$$T = A \cup B \cup C, \quad (1)$$

тогда как на практике при решении конкретных задач таблицы T состоят из сочетаний (компонентов):

$$T = A_1 \cup A_2 \cup B \cup C, \quad (2)$$

$$T = A_{1(2)} \cup B \cup C, \quad (3)$$

$$T = A_{1(2)} \cup B_{1(2)} \cup C, \quad (4)$$

$$T = B \cup C, \quad (5)$$

$$T = B_{1(2)} \cup C, \quad (6)$$

где $1(2)$ означает возможность наличия одного из индексов при A и B .

Если $B_{1(2)}$ в (6) соответствует одному домену, имеем одностолбцовую (однотрочную) таблицу, где содержимое $B_{1(2)}$ описывает

запрос на множество показателей; если же содержимое $B_{i(2)}$ описывает запрос на один показатель, то имеем случай таблицы размером $I \times I$ (бестабличный запрос).

Нетрудно заметить, что (2) – (4) могут быть сведены к (5), так как содержимое $A_1 \cup A_2$ можно отнести к одному верхнему домену части B_i , ширина которого равна ширине B_i . В связи с этим все рассуждения приводятся для таблицы, заданной формулой (5).

Несмотря на успехи, достигнутые в области создания баз данных (БД) и средств теледоступа, в АСУ еще доминирует позадачный подход, когда для каждой Т составляются специальные программы, определяющие и при необходимости обрабатывающие первичные показатели с целью вычисления значений агрегируемых показателей заданной Т. При этом роль БД сводится к хранению и выдаче значений первичных показателей, а роль средств теледоступа – к ускорению процесса подготовки и ввода-вывода программ и данных.

Жесткая зависимость программ решения задач АСУ от структуры Т и содержания $B \subset T$, а отсюда – необходимость привлечения специалиста программиста при всевозможных изменениях Т, значительный коэффициент дублирования программных средств и памяти ЭВМ, привязка пользователей системы обработки данных (СОД) к фиксированному набору таблиц представления задач являются причинами, сдерживающими внедрение диалоговых методов решения задач непосредственно пользователями непрограммистами.

В последние годы предложены новые подходы к решению задачи автоматизации процесса заполнения Т. Так, в [1] рассматриваются дополнительные средства к БД фактографического типа, построенной согласно [2], которые позволяют пользователям непрограммистам заданием одних только запросов на формализованном языке, близком к естественному, формируемых ими на базе $B \subset T$, получить значение первичных и агрегируемых показателей данной предметной области.

Решение задач в режиме взаимодействия пользователей с СОД, организованными в соответствии с [1; 2], оставляет за пользователями право на выполнение процедуры "расшивки" требуемых выходных таблиц (ТВТ) с целью составления из $B \subset T$ запросов L на формализованном языке, которые однозначно описывают все показатели, содержащиеся в Т. Это усложняет процесс формирования и ввода запросов (т.е. задач) в ЭВМ как при пакетной, так и при диалоговой обработке данных. К тому же процедуры формирования и вывода ТВТ сильно усложняют программы обработки Т.

Таким образом, возникает необходимость разработки средств автоматического преобразования $(B \subset T) \rightarrow L$, в результате чего стало бы возможным инвариантное по отношению к Т решение задач извест-

ными средствами с одновременным исключением процедур формирования выходных таблиц.

Постановка задачи

На практике используемые таблицы часто содержат в доменах B лишние подграфы, вспомогательную информацию в виде аббревиатур, комментариев и т.п. Легко заметить, что все таблицы могут быть приведены к нормализованному виду.

Определение 1. ТБТ называется нормализованной, если в ней: а) отсутствуют домены $D_k \in B_{I(2)}$ одинаковой ширины (высоты), охваченные общими вертикальными (горизонтальными) линиями; б) каждый домен содержит только комбинации элементов заданного тезауруса рассматриваемой предметной области.

П о с т а н о в к а з а д а ч и 1. На формализованном языке, близком к естественному, задан тезаурус рассматриваемой предметной области $[I]$, в котором выделяются множества наименований подлежащих S , признаков X и значений признаков Z показателей, элементы которых достаточны для однозначного описания всего множества L первичных и агрегируемых показателей.

Пользователи задают незаполненные нормализованные ТБТ всевозможных конфигураций T , такие, где каждый запрос, формируемый содержанием соответствующих доменов частей $B_{I(2)}$, содержит не более одного элемента множества S .

Необходимо установить алгоритм, предусматривающий автоматическое преобразование $(B \subset T) \rightarrow L$, в результате чего станет возможным преобразование $B \rightarrow C$ без участия пользователя СҚД.

П о с т а н о в к а з а д а ч и 2. Заданы условия задачи 1 с тем лишь отличием, что ТБТ заполнены значениями первичных показателей.

Необходимо установить алгоритм, предусматривающий автоматическое преобразование $B \rightarrow L$ и запись чисел множества C в БД без участия пользователя СҚД.

Решение задачи 2 сводится в основном к решению задачи 1, но после этапа синтеза запросов не выполняется процедура нахождения значений показателей, а производится запись в БД значений соответствующих показателей.

Описывая алгоритм, будем исходить из того, что значения первичных (неагрегируемых) показателей хранятся в БД фактографического типа, построенной согласно [2] или [4]. В целом же сам алгоритм не зависит от типа используемой БД.

Определения и условные обозначения

Прежде чем приступить к рассмотрению самого алгоритма, введем некоторые новые понятия и обозначения.

Для просмотра и анализа T с целью формирования запросов (предложений), однозначно описывающих показатели $C \subset T$, воспользуемся известным из теории распознавания зрительных образов способом сканирования входного образа (в данном случае – нормализованной ТВТ) и некоторыми положениями работы [3]. Причем, можно применять как построчную, так и контурную развертку (т.е. развертку по доменам).

Вначале рассмотрим алгоритмы синтеза запросов с использованием построчной развертки T , а затем отметим особенности синтеза запросов для случая контурной развертки.

Введем обозначения:

T_{ij} – заданная ТВТ (вся совокупность входных символов таблицы), где i, j – соответственно номера строк и столбцов развертки ТВТ;

k, p – порядковый номер доменов в процессе сканирования T_{ij} слева направо: D_k соответствует доменам i -й (текущей) строки, а D_p – доменам $(i-1)$ -й (предыдущей) строки;

h – номер вертикальной линии, пересекающей текущую сканируемую строку;

n_k, u_k – значения j , соответствующие началу и концу текста k -го домена текущей строки;

$y(2k-1), y(2k)$ и $\hat{y}(2p-1), \hat{y}(2p)$ – значения j , указывающие границы доменов соответственно по текущей и предыдущей строкам (в дальнейшем – вход и выход доменов);

$\varphi_k(\varphi_p)$ – признак наличия (1) или отсутствия (0) горизонтальной линии в k -м (p -м) домене текущей (предыдущей) строки;

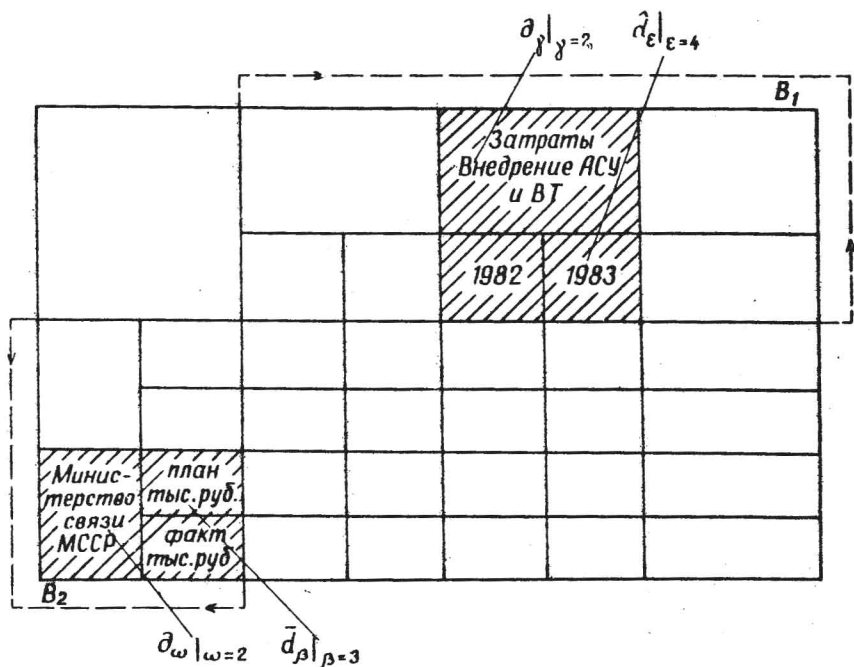
F_k – текст k -го домена текущей строки; после обработки информации текущей строки – текст k -го домена по i -ю строку включительно;

$\hat{F}_p$ – текст p -го домена по $(i-1)$ -ю строку включительно;

w – значение j , соответствующее левой границе $C \subset T$;

$w = y(2k-1) \Big|_{\substack{k=1 \\ D_k \in B}}$.

Для доменов D_k и D_p введем понятие перекрытия, определяемое как $\sigma_{kp} = \text{sign} [y(2k) - \hat{y}(2p-1)] \cdot \text{sign} [\hat{y}(2p) - y(2k-1)]$. При этом, если $\sigma_{kp} = 1$, следует говорить о наличии перекрытия между доменами k и p и их содержимое относить к одному и тому же запросу (предложению); если $\sigma_{kp} \neq 1$, то домены k и p не имеют пе-



Р и с. 1. Пример нормализованной ТБТ

рекрытия, и их содержимое относится к разным запросам.

Определение 2. Домен $D_k \in B$, внутри которого $\varphi_k \neq 1$, назовем доменом I-го рода (d_k).

Определение 3. Домены I-го рода, граничащие с $C \in T$, назовем базовыми вертикальными $\bar{d}_\epsilon \in B_1$ и базовыми горизонтальными $\bar{d}_\beta \in B_2$.

Определение 4. Доменом 2-го рода d_γ (d_ω) назовем составной домен, образованный одним из крайних сверху (слева) доменов I-го рода части $B_{(2)}$ и множеством всех доменов I-го рода, которыми он перекрывается.

Домен d_k , имеющий высоту (ширину), равную высоте (ширине) части $B_1 \in T$ ($B_2 \in T$), также отнесем к доменам 2-го рода.

Значения индексов k, p, γ, ϵ растут слева направо, а β, ω — сверху вниз (рис.1).

Исходя из сформулированных определений, количество запрашиваемых показателей ТБТ m определяем произведением числа базовых вертикальных ϵ_0 на число базовых горизонтальных β_0 доменов: $m = \epsilon_0 \cdot \beta_0$.

Алгоритм синтеза запросов по заданной ТВТ

При описании рассматриваемого алгоритма приняты следующие допущения, которые не носят принципиального характера, но упрощают его машинную реализацию: а) между терминами (словами) – элементами множеств S, X, Z , расположенными в одной строке, разрешается только один пробел; б) наименования подлежащих (т.е. элементы множества S) в заданной ТВТ включаются либо в крайние сверху домены I-го рода части B_1 , либо в крайние слева домены I-го рода части B_2 .

Алгоритм основан на следующем утверждении: для однозначного определения содержащихся в ТВТ запросов достаточно один раз сканировать $B_{1(2)}$ сверху вниз и слева направо и анализировать выделенные в процессе сканирования значения установленных признаков только для i -й и $(i-1)$ -й строк, а также содержимое просмотренных горизонтальных базовых доменов $\bar{d}_p$ заданного сканируемого домена 2-го рода для $\bar{d}_\omega \in B_2$.

Необходимость хранения содержимого просмотренных горизонтальных базовых доменов заданного сканируемого домена 2-го рода объясняется неодинаковым размещением доменов частей $B_{1(2)}$ таблицы ($D_k \in B_1$ размещены вертикально, а $D_k \in B_2$ – горизонтально). Если текстовую информацию $D_k \in B_2$ разместить по столбцам (снизу вверх), как это делается в ряде случаев, то необходимость в хранении такой информации отпадает: остается лишь поменять местами значения параметров i и j при анализе частей $B_{1(2)}$ таблицы.

В реализованной версии данного алгоритма нахождение значений ранее введенных параметров осуществляется путем обработки литерных строк $B \in T$ согласно следующим соотношениям:

$$\begin{aligned}
 h &= \begin{cases} h+1, & \text{если } T_{i,j} = T_{i-1,j} \vee T_{i+1,j} = " | " ; \\ h & \text{в противном случае,} \end{cases} \\
 \varphi_k &= \begin{cases} 1, & \text{если } T_{i,j} = T_{i,j+1} \vee T_{i,j-1} = " - " ; \\ 0 & \text{в противном случае} \end{cases} \\
 k &= \begin{cases} k+1, & \text{если } T_{i,j} = " | ", h > 1 |_{w \neq 0}, \\ & \text{или } T_{i,j} = " | ", T_{i,j-1} \neq " | " ; \\ k & \text{в противном случае} \end{cases} \\
 n_k &= \begin{cases} j, & \text{если } T_{i,j-2} = T_{i,j-1} = " - ", T_{i,j} \neq " - " \vee " | " \vee " - " ; \\ & \text{или } T_{i,j-1} = " | ", T_{i,j} \neq " - " \vee " - " ; \\ n_k & \text{в противном случае} \end{cases}
 \end{aligned}$$

$$u_k = \begin{cases} J, & \text{если } T_{i,j+1} = T_{i,j+2} = \text{"—"}, T_{i,j} \neq \text{"—"}, \text{"—"}, \text{"—"}, \text{"—"} \\ & \text{или } T_{i,j} \neq \text{"—"}, \text{"—"}, \text{"—"}, \text{"—"}, T_{i,j+1} = \text{"—"} \\ u_k & \text{в противном случае} \end{cases};$$

$$y(2k-1) = \begin{cases} J, & \text{если } y(2k-1) = 0, k|_{T_{i,j-1}} = a, k|_{T_{i,j}} = a+1 \\ 0 & \text{в противном случае} \end{cases};$$

$$y(2k) = \begin{cases} J, & \text{если } y(2k-1) \neq 0, k|_{T_{i,j-1}} = a, k|_{T_{i,j}} = a+1 \\ 0 & \text{в противном случае} \end{cases};$$

$$w = y(2k-1)|_{k=1}^{h=2} = y(1);$$

$$F_k = F_k \parallel T_{i,j}, \\ j = \overline{n_k, u_k}$$

где „|“, „—“ — элементы соответственно вертикальных и горизонтальных линий; „—“ — пробелы; $\vee, \parallel$ — операции соответственно дизъюнкции (ИЛИ) и сцепления.

Алгоритм синтеза запросов

1. Ввод $T_{i,j} (i = \overline{1, I_0}, j = \overline{1, J_0})$. Установление исходных значений рассмотренных ранее параметров.

2. Выделение и анализ признаков доменов $D_k \in B_1$.

2.1. $i = i+1, j = j+1$. Поиск первого элемента $T_{i,j} \neq \text{"—"}, \text{"—"}, \text{"—"}, \text{"—"}$. Если найденный элемент $T_{i,j} = \text{"—"}, \text{"—"}, \text{"—"}, \text{"—"}$, переход к п.2.3.

2.2. Случай $A_{1(2)} \neq \emptyset$. $\hat{F}_1 = \hat{F}_1 \parallel T_{i,j}, T_{i,j} \in A_{1(2)}$. В процессе сцепления элементов $T_{i,j}$ исключаются двойные пробелы. Переход к п.2.1.

2.3. $i = i+1, j = 1$. Сканируется строка i , пока $T_{i,j}|_{h=1} = \text{"—"}, \text{"—"}, \text{"—"}, \text{"—"}$. Если $T_{i,j+1} = T_{i,j+2} = \text{"—"}, \text{"—"}, \text{"—"}, \text{"—"}$ & $w \neq 0$, переход к п.2.7.

2.4. Находится $T_{i,j}|_{h=2} = \text{"—"}, \text{"—"}, \text{"—"}, \text{"—"}$: Далее в процессе сканирования i -й строки фиксируются значения $y(2k-1), y(2k), n_k, u_k, \varphi_k$ и F_k для всех k . При этом, если $w = 0$, то $w = y(1)$.

2.5. $\forall k$ и p определяются ϕ_{kp} . Если $\phi_{kp} = 1$, то $F_k = \hat{F}_p \parallel F_k$.

2.6. Выполняются действия

$$\left. \begin{aligned} \hat{y}(2a-1) &= y(2a-1), \hat{y}(2a) = y(2a), \hat{\varphi}_a = \varphi_a, \hat{F}_a = F_a \\ y(2a-1) &= y(2a) = \varphi_a = 0; F_a = \text{"—"} \end{aligned} \right\} a = \overline{1, k}$$

$$p = k, k = 0. \text{ Переход к п.2.3.}$$

2.7. $e = p$.

$$\left. \begin{aligned} \varphi_p &= \hat{F}_p, \hat{F}_p = \text{"—"} \\ y(2p-1) &= y(2p) = \varphi_p = 0 \end{aligned} \right\} p = \overline{1, e}, \\ p = 0$$

где $\varphi_p (p = \overline{1, \epsilon})$ - приведенное содержимое базовых вертикальных доменов; " - строка нулевой длины для переменных строкового типа.

3. Выделение и анализ признаков доменов $D_k \in B_2$.

3.1. $i = i + 1, j = 1$. Находится $T_{ij} |_{h=1} = "1"$. Если $T_{ij} \neq "1"$ для $j = \overline{1, j_0}$, переход к п.4.

3.2. Сканируется строка до $j = \omega$, фиксируя при этом значения $y(2k-1), y(2k), n_k, u_k, \varphi_k$ и F_k для всех выделенных k .

3.3. $\forall k$ и p определяются σ_{kp} . Если $\sigma_{kp} = 1$, то $F_k = \hat{F}_p \parallel F_k$.

3.4. Анализируется значение признаков домена, для которого $y(2k) = \omega$. Если $\varphi_k \neq 1$, переход к п.3.8.

3.5. $\beta = \beta + 1$,

$\lambda_\beta = k - 1$,

$G_\beta = F_k$,

$$q = \begin{cases} \beta, & \text{если } q = 0 \\ q & \text{в противном случае;} \end{cases}$$

где λ_β - количество доменов, расположенных левее k -го домена $D_k \in B_2$; G_β - содержимое β -го горизонтального базового домена; q - номер первого горизонтального базового домена рассматриваемого домена 2-го рода.

3.6. Если $\prod_{t=1}^k \varphi_t \neq 1$, то $F_k = "$, $k = k - 1$. Переход к п.3.8.

3.7. $G_r = \left\{ \parallel_{t=\overline{1, \lambda_r}} F_t \right\} \parallel \theta_r, r = \overline{q, \beta}$;

$$\left. \begin{aligned} y(2a-1) = \hat{y}(2a-1) = \hat{y}(2a) = y(2a) = \varphi_a = 0 \\ \hat{F}_a = \hat{F}_a = " \\ p = k = 0 \\ q = \beta + 1. \end{aligned} \right\} a = \overline{1, k}$$

Переход к п.3.1.

$$\left. \begin{aligned} \hat{y}(2a-1) = y(2a-1) \\ \hat{y}(2a) = y(2a) \\ \hat{\varphi}_a = \varphi_a \\ \hat{F}_a = F_a \\ y(2a-1) = y(2a) = \varphi_a = 0 \\ F_a = " \end{aligned} \right\} a = \overline{1, k}$$

3.9. Переход к п.3.1.

4. Формирование результирующих запросов ТВТ:

$$L_{ij} = \begin{cases} \varphi_j \parallel \theta_i, & \text{если } s_i \in S \text{ включены в } B_1 \\ \theta_i \parallel \varphi_j, & \text{если } s_i \in S \text{ включены в } B_2 \end{cases}$$

где $i = \overline{1, \beta}, j = \overline{1, \epsilon}$; L_{ij} - результирующее семантическое описание (i, j) -го показателя (запроса) ТБГ.

Нахождение значений c_{ij} (в БД или в результате обработки значений других показателей) требует, однако, перехода от семантического к кодовому описанию $L_{ij} \in L$. Это предполагает разбиение предложения L_{ij} на соответствующие элементы множеств заданного тезауруса.

В общем случае предложение L_{ij} представлено совокупностью элементов $L_{ij} = \{s_i, \bigcup_{p \in \Omega_z(i, j)} z_p, \bigcup_{t \in \Omega_x(i, j)} x_t\}$ или символов $L_{ij} = \{l_k; k = \overline{1, \xi}\}$,

где l, n, t - конкретные для (i, j) -го показателя номера элементов множеств S, Z, X ; l_k - символы формализованного языка запросов; ξ - длина запроса. Заметим, что для любых $L_{ij} \Omega_z(i, j) \neq \emptyset$, тогда как иногда $\Omega_x(i, j) = \emptyset$.

Разбиение L_{ij} на элементы множеств S, Z, X осуществляется следующим образом.

Предварительно массивы S, Z, X ранжируются в порядке убывания длин их элементов (количества образующих их символов); значения длин хранятся в соответствующих элементах массивов длин данных множеств v_s, v_z, v_x .

С учетом того, что s_i всегда находится в начале L_{ij} , а $\Omega_z(i, j) \neq \emptyset$, наборы символов l_k сравниваются до первого совпадения последовательно с элементами S, Z, X .

Определение кодов L_{ij} предусматривает выполнение следующих шагов:

1. $E = S = \{s_k\}, k=0, p=1$.
2. $k = k+1$.
3. Если $k > k_{max}$, переход к п.6.
4. Если $e_k = \left\{ \bigcup_{n=p, v_z^t} l_n \right\}$, то

$$p = v_z^t + 1,$$

$$L_{ij} = \{\bar{\theta} \parallel \theta_k \parallel (\bigcup_{m=p, q} l_m)\},$$

где $\bar{\theta}$ - коды элементов, найденные на предыдущих шагах; θ_k - код элемента e_k . Переход к п.8.

5. Переход к п.2.

6. Если $E = S$, то $E = Z = \{z_k\}, k=0$, переход к п.2.

7. Если $E = Z$, то $E = X = \{x_k\}, k=0$, переход к п.2.

8. Если $p \neq q$, то $E = Z = \{z_k\}, k=0$, переход к п.2.

9. Конец.

Реализация алгоритма определения содержимого L_{ij} способом контурной развертки основывается также на выделении значений установленных признаков согласно формуле (7) за исключением e_{kr} . Отличие заключается лишь в том, что значения указанных признаков определяются только для одного, сканируемого домена; необходимо дополнительно хранить полученное содержимое всех доменов I-го рода заданного домена 2-го рода, пока не закончится его сканирование.

Таким образом, реализация данного алгоритма позволяет наряду со средствами, рассмотренными в [1; 2], использовать следующую технологию обработки задач АСУ: а) составление на экране терминала ТБТ; б) приведение ТБТ к нормализованному виду; в) запуск программы синтеза запросов; г) запуск программы определения значений показателей; д) вывод на экран терминала значений показателей c_{ij} ; е) добавление (или изменение) при необходимости информации в $A_{1(2)}$ и/или $B_{1(2)}$; ж) распечатка результирующей ТБТ на АЦПУ.

Составлены и апробированы PL/I-программы, реализующие обе версии данного алгоритма, для различного рода ТБТ как в пакетном, так и в диалоговом режиме. Апробирована работа данного алгоритма совместно с известными средствами при решении в пакетном режиме конкретной задачи для одного из отделов Госплана МССР.

Л и т е р а т у р а

1. Б у ж о р П. Ф. Обработка информации в АСУ с применением баз данных фактографического типа и семантической идентификации информации. Деп. рукопись. МолдНИИТИ Госплана МССР. - Библиографический указатель "Депонированные рукописи", 1983, № 1.

2. К о р н я С. К. Инструментарий построения безызбыточных баз данных. - В кн.: Информационные языки и базы данных в планировании. М., ЦРМИ АН СССР, 1980.

3. Н е д е л ь ч у к В. П. Вопросы идентификации фигур с использованием оптоэлектронных параллельных преобразователей изображения. - Сб. трудов МАИ, вып. 461. М., 1978.

4. Н е д е л ь ч у к В. П. Использование методов ассоциативной и прямой адресации для организации фактографических БД АСУ. - В кн.: Методология АСУ. Кишинев, Штиинца, 1982.