

ВИКОРИСТАННЯ КЛАСИФІКАЦІЙНИХ ПРАВИЛ ДЛЯ ЗМЕНШЕННЯ НЕВИЗНАЧЕНОСТІ У СХОВИЩАХ ДАНИХ, ПОБУДОВАНИХ НА ОСНОВІ РЕЛЯЦІЙНОЇ МОДЕЛІ

© Шаховська Н. Б., Кісь Я.П., 2005

Описано методи побудови класифікаційних правил для усунення невизначеностей. Запропоновано алгоритми для розв'язання основних задач.

The methods of uncertain eliminate are described. The main algorithms of design are described.

Вступ. Як відомо, до сховищ даних (СД) часто формують запити з нечітко заданими параметрами, поданими у вигляді інтервалів, лінгвістичних змінних, ступенів довіри тощо. Крім того, усе частіше виникає проблема збереження у відношеннях СД, побудованого на основі реляційної моделі, інформації про об'єкти, які є нечіткими. Прикладом таких об'єктів є різноманітні класи, причому класифікаційні ознаки можуть бути як чіткими (наприклад, загальноприйнята вікова градація *Молодий – до 28 років, Зрілий – від 28 до 40 років, Старий – 40 років і старший*), так і заданими певним розподілом.

Найчастіше нечіткість може з'являтися у переглядах, отриманих у результаті запиту до СД із параметрами, поданими у вигляді нечітких величин. Методи та правила здійснення вибірки, а також шкали для порівняння добре відомі.

Типовими предметними галузями, у яких постає задача опрацювання (усунення) невизначених та нечітко заданих значень, є: задачі планування, консалтингові компанії, соціологічні опитування, історичні дослідження тощо.

Основними проблемами, які виникають в задачах аналізу та структурування даних, є проблеми створення класів та віднесення до них об'єктів, інформація про які щойно надійшла у сховище даних.

Пропонуються алгоритми класифікації та класифікування об'єктів, інформація про які зберігається у сховищі даних.

Огляд сучасних досліджень та виділення невирішених проблем. Оскільки розглядають СД, побудоване на основі реляційної моделі, то у нього можуть бути введені ті самі типи невизначеності, що й у відношення реляційної бази даних [2]:

1. Значення невідоме (відсутнє).
2. Неповнота інформації.
3. Нечіткість (використання розподілу для встановлення істинності знань).
4. Неточність (стосується числових даних).
5. Недетермінованість процедур виведення рішень (випадковість).
6. Ненадійність даних.
7. Багатозначність інтерпретацій.
8. Лінгвістична невизначеність.
 - а) Невизначеність значення слова.
 - б) Невизначеність змісту речення.

У відношення реляційної бази даних уведено перші три типи невизначеності, а до решти типів нема чіткого означення та методів їхнього розпізнавання.

Тому дамо формальне визначення типам невизначеності та проаналізуємо випадки їхньої появи у відношеннях сховища даних.

Відсутність інформації трапляється на рівні значення атрибута і означає, що значення притаманне об'єкту, але невідоме

$$t = \{A, unk\},$$

де t – кортеж, unk – атрибут із відсутнім значенням, A – решту значень атрибутів кортежу, $unk \cup A = t, unk \cap A = \emptyset$.

Прикладами появи такої невизначеності є відсутність ідентифікаційного номера працівника фірми, затримка у надходженні накладних про продаж (невідома кількість проданого товару) тощо.

Неповнота є станом кортежу, у якому є підмножина відсутніх значень атрибутів. Якщо ця підмножина порожня, то ми отримуємо традиційний кортеж. Відсутність інформації є також частковим випадком неповноти інформації, коли кількість невідомих значень атрибутів кортежу дорівнює 1.

$$t = \{A, \{unk\}\},$$
$$|unk| < |A|$$

Прикладом неповноти інформації є відсутність ідентифікаційного номера та номера паспорта.

Невизначеності типів 3 – 8 класифікують у [2] як неоднозначність даних і переважно з'являються на рівні кортежу або підмножини значень атрибутів, із яких формується кортеж.

Нечіткість виникає через неповне вивчення або неоднозначне відображення характеристик сутності у відношенні. Моделюється за допомогою додавання до схеми відношення додаткового атрибута (атрибутів), значення яких містять рівень впевненості у істинності підмножини значень неключових атрибутів

$$t = \{A, unk_1, \dots, unk_n\},$$
$$A \in K, A',$$
$$1 \leq n \leq |A'|$$

де K – множина значень ключів, A' – підмножина відомих значень неключових атрибутів. Рівень впевненості може позначатися за допомогою числової шкали, лінгвістичних оцінок, нечіткої величини тощо.

Приклад кортежу з нечіткістю

Дата доку-мента	Тип опера-ції	Код замовл.	Код менеджер.	Код підро-зділу	Код клієнта	Код угоди	Код товару	К-ть	Unk1	Ціна	Unk2	Сума	Unk3
2.2.05	прихід	04	Менджер1	Склад1	Клієнт1	Угода1	Сист блок	2	0.4	22		44	
2.3.05	витрата	01	Менджер1	Склад1	Клієнт1	Угода1	монітор	1		23	0.8	23	
3.6.05	прихід	03	Менджер1	Склад1	Клієнт2	Угода1	клавіатура	4		2		8	0.6

Тут значення атрибута $unk1$ відображає ступінь довіри до значення атрибута *кількість*, $unk2$ – атрибутів *первинна ціна*, *вторинна ціна*, $unk3$ – атрибутів *первинна сума*, *вторинна сума*.

Неточність отримується внаслідок застосування математичних операцій над числовими даними (до них зараховують невизначеність, яка виникає внаслідок роботи з інтервальними величинами). Цей тип невизначеності не моделюється за допомогою якихось додаткових засобів і

важко виявляється на рівні кортежів. Нині немає засобів запобігання появи такого типу невизначеності.

Наведемо приклад виникнення неточності. Нехай є відношення *Плани випуску продукції*, у якому прогнозована кількість продукції моделюється за допомогою двох атрибутів, що подають нижню межу та верхню межу кількості відповідно. Ставиться задача визначення залежності кількості продукції від пори року. Тоді застосування математичних операцій над значеннями інтервалу кількості призведе до розширення інтервалу [7, с. 878], що, відповідно, зменшить точність прогнозу.

Недетермінованість процедур виведення рішень (випадковість) виникає, коли необхідно зберігати проміжні або кінцеві результати процедур виведення або прийняття рішень. Оскільки отримані під час виведення оцінки є неточними, такий вид невизначеності вважається фізичним і також не може бути попереджений.

Ненадійність є типом невизначеності, який вважають однією із характеристик об'єкта. Наприклад, у відношенні *Замовлення* значення атрибута *Ймовірність виконання* визначають на основі попереднього аналізу характеристик замовлення та зберігають у відношенні для інформативності. Хоча сама природа цієї характеристики є невизначеною, у відношенні як її домен використовують традиційну числову шкалу (наприклад, від 0 до 1) та виконують над її значеннями традиційні математичні операції.

Багатозначність інтерпретації є одним із джерел виникнення суперечностей. Такий тип невизначеності виникає найчастіше стосовно фактів через отримання інформації із різних джерел і неможливість визначення істинної. Для представлення цього типу невизначеності до схеми відношення додають додатковий атрибут, який містить ступінь впевненості в істинності даних кортежу. Від типу нечіткості відрізняється тим, що вводиться на рівні відношення.

Приклад появи багатозначності інтерпретації:

Нехай у документі про торгівлю, який надійшов у вигляді твердої копії, погано видно кількість одиниць товару. Тоді у відношенні з'являються два кортежі, які відрізняються між собою значеннями виділених жирним шрифтом атрибутів.

код	Дата доку-мен-та	Тип опе-ра-ції	Код замов-лення	Код менед-жера	Код підроз-ділу	Код клієнта	Код угоди	Код товару	К-ть	Unkl	Ціна	Су-ма
01	2.2.0 5	при-хід	04	Мене-джер1	Склад1	Клієнт 1	Угода1	Сист блок	6	0.4	22	44
02	2.2.0 5	при-хід	04	Мене-джер1	Склад1	Клієнт 1	Угода1	Сист блок	8	0.6	22	44

Лінгвістична невизначеність пов'язана із використанням природної мови для представлення знань, які мають якісний характер, і може виникати внаслідок нерозуміння (незнання) значення слова або нерозуміння змісту речення. Такий тип невизначеності трапляється у системах обробки текстової інформації (системи автоматизованого перекладу, системи для самонавчання тощо).

Розглянуті типи невизначеностей можуть накладатись один на одного або бути джерелом появи один одного.

Нині розроблено методи усунення відсутніх, неповних та нечітких даних [2,4]. Тому необхідно розробити методи, які б працювали з усіма типами невизначеності.

Постановка задачі. Вхідними даними для класифікування (зарахування до класу) є множина значень цільових атрибутів. *Цільовими атрибутами (X)* назвемо атрибути, які використовують для аналізу даних, і згідно із значеннями яких здійснюють розбиття на класи. До

цільових атрибутів, передовсім, належать усі атрибути, які входять до множини ключів. До цільових атрибутів зарахуємо усі атрибути, що входять у множину лівих частин функціональних залежностей (крім первинних ключів), а також ті атрибути, які будуть впливати на ступінь довіри до отриманого результату аналізу. Крім того, для конкретної предметної галузі за допомогою експертного опитування визначається додаткова підмножина атрибутів, які вважатимуть цільовими для аналізу. Наприклад, для задачі соціологічного опитування такими атрибутами є вік, освіта, матеріальний стан тощо.

Атрибути, над якими виконують операції агрегації та порівняння, назовемо *критичними* Y (подають результати аналізу). *Критичними* атрибутами є атрибути, які містять числові дані, невизначеності, подані у довільному вигляді, та праві частини функціональних залежностей. До них також належать атрибути, що містять назви класів (*мітки*) [4].

Класом назовемо підмножину кортежів, для яких значення по множині критичних атрибутів є однаковими

$$cl = \sigma_r(X = x, Y = y).$$

Міткою класу назовемо лінгвістичну змінну або типову характеристику об'єктів, яка є значеннями підмножини атрибутів Y і позначає об'єкти зі спільними (подібними зі ступенем s) значеннями підмножини атрибутів X . Домени атрибутів, що належать до підмножини Y , $y \in \text{dom}(Y) = \pi_Y(r)$, обов'язково повинні містити скінченну та наперед відому множину значень.

Для спрощення задачі вважатимемо, що класи є визначимими, і їхні характеристики (тобто назви та правила, за якими об'єкт вважається представником цього класу) зберігаються у сховищі даних.

Ставимо задачу: маємо деяке відношення r зі схемою R . Для усунення невизначеності на основі відношення r необхідно побудувати множину класифікаційних правил виду $cl(X \rightarrow Y)$, де $X, Y \subset R$, $X \cap Y = \emptyset$; X – підмножина цільових атрибутів, Y – підмножина критичних атрибутів [1].

До класу зараховують на основі визначення підмножини значень цільових атрибутів. Наприклад, для класу “Термінове замовлення” значення цільових атрибутів мають задовольняти умови: дата закінчення терміну – $[\text{DateDiff}(\text{"d"}, \text{Now}(), 30), \text{Now}())$ (тут $\text{Now}()$ – сьогоднішня дата, операція DateDiff віднімає від параметра №2 кількість одиниць, вказаних у параметрі №3 типу параметр №1. Результатом операції DateDiff буде дата, менша на 30 днів за теперішню дату), сума – більша від 0, тип операції – прихід.

Оскільки важко отримати повну інформацію про об'єкти предметної галузі, то можуть бути визначені не всі цільові атрибути. Тому для кожного класу визначають значення *межі* – величини у межах одиничного інтервалу, яка позначає мінімальний ступінь довіри до об'єкта, за яким об'єкт може бути класифікований як представник цього класу. Ступінь довіри s до об'єкта визначають як кількість цільових атрибутів із визначеними значеннями до усіх визначених цільових атрибутів цього класу (чим більше відомо про об'єкт, тим вищим буде ступінь довіри).

$$s = \sum \begin{cases} 0, & cr_i \notin cl \\ 1, & cr_i \in cl \end{cases}$$

де cr_i – значення по множині цільових атрибутів.

Віднесення до класу можна розглядати як один із способів усунення невизначеності, адже під час класифікування здійснюється заповнення порожнього значення атрибута, який містить значення назви класу. Крім того, класифікаційні правила можна вважати нечіткими (наближеними) функціональними залежностями.

У базі даних підтримується нечітка функціональна залежність

$$e(X \rightarrow A),$$

якщо співвідношення кортежів, на яких виконується ця функціональна залежність, до кортежів, на яких вона не виконується, не менше ніж s , де s – значення межі пропускання, визначене на основі експертного опитування [3]. Зрозуміло, значення s – не менше ніж значення межі класу.

$$e(X \rightarrow Y) : \frac{COUNT(X = x, Y = y)}{COUNT(X = x, Y \neq y)} \geq s_{cl}. \quad (1)$$

Значення межі пропускання змінюється у межах $[0, 1]$.

Звідси випливає, що алгоритми усунення невизначеностей за допомогою функціональних залежностей можна застосувати для класифікування об'єктів.

Як приклад, для розглянутої вище предметної галузі планування замовлень, у сховищі даних існує класифікаційне правило *Дата, Сума, Тип операції* $\rightarrow$ *Тип замовлення*.

Застосування класифікаційних правил для зменшення невизначеності порівняно з використанням лише традиційних функціональних залежностей дають змогу:

- Моделювати особливості конкретної предметної галузі.
- Значно розширити коло невизначеностей, які можуть бути усунені (нагадаємо, що за допомогою функціональних залежностей усувалися лише невідомі значення).
- Використовуючи знання (досвід) експерта, видобувати нові знання.

Основний матеріал. Для того, щоб мати можливість класифікувати об'єкти, необхідно побудувати правила класифікації. Взагалі у сховищі даних може зберігатися інформація про декілька типів класів, і для кожного типу класу є своя підмножина правил. Одне й те саме правило може застосовуватись для визначення кількох типів класів.

Правила можуть генеруватись двома способами:

- на основі аналізу характеристик класів;
- на основі чинних правил.

А. Породження класифікаційних правил на основі аналізу характеристик класу

У класифікаційні правила, передовсім, будуть входити атрибути, що визначають з функціональних залежностей. Решту атрибутів визначають на основі аналізу характеристик класів.

Послідовність кроків:

- 1) Кортежі відношення групують за назвами класів.
- 2) Вибираємо перший із цільових атрибутів $i=1$. Утворюємо послідовність групування $G = \emptyset$.
- 3) Якщо $i \geq |R \setminus \{Y\}|$, перейти на крок 8.
- 4) Всередині групи за назвами класу проходить групування за кожним цільовим атрибутом. Порядок групування визначається множиною G , після вичерпання елементів якої відбувається за рештою цільових атрибутів.
- 5) Визначаємо значення e_i за формулою (1). Якщо $e_i \rightarrow \infty$ ($COUNT(X_i = x_i, Y = y) = 0$), цільовий атрибут X_i безпосередньо впливає на Y і ми знайшли частину функціональної залежності. Тоді $i=i+1$ та $G = G \cup X_{i+1}$. Перейти на крок 3.
- 6) Якщо $e_i > s_{cl}$ і $\forall k (X_k = x_k), (X_i = x_i) \rightarrow (Y = y)$ і $e(X_1, \dots, X_k, X_i) > s_{cl}$, $k = 1, |G|$, то $i=i+1$ та $G = G \cup X_{i+1}$. Якщо $e(X_1, \dots, X_k, X_i) < s_{cl}$, то X_i включатимемо у нове правило. Для цього назви усіх атрибутів з G зберігаємо у новому кортежі відношення правил, $A = e$ – довіра до сформованого правила, $G = \{X_i\}$

7) $i=i+1$. Перейти на крок 3.

8) Якщо $G = \emptyset$, класифікаційних правил не знайдено.

В. Породження класифікаційних правил на основі відомих

Відомо, що за допомогою логічного виведення можна отримувати нові правила, задані у вигляді “якщо, то” [5]. Класифікаційні правила мають саме такий вигляд. У [4] введено поняття наближеної функціональної залежності, яка має встановлене значення ступеня довіри. Якщо класифікаційні правила розглядати як нечіткі (наближені) функціональні залежності зі ступенем довіри S , то до них можна застосувати основні аксіоми виведення [8]. Оскільки ми маємо справу з дискретними величинами, то для опрацювання ступенів довіри до правил використаємо апарат багатозначної логіки Лукасевича (проаналізовано у [5]. Тоді використання правил виведення та застосування логічних операції “і” для правих частин та “або” для лівих дасть можливість генерувати нові правила на основі чинних та автоматично визначати ступені довіри до них (які можуть бути перевірені експериментально).

Оскільки доведено, що для наближених функціональних залежностей є істинними аксіоми виведення [4], то звідси випливає, що у базі знань потрібно зберігати лише мінімальне покриття наближених функціональних залежностей (тобто класифікаційних правил), а решту можна вивести на основі їхніх комбінацій з використанням операцій багатозначної логіки (ґрунтується на мінімаксному підході) та аксіом виведення.

С. Алгоритм усунення невизначеності за допомогою класифікаційних правил

Розглянемо один із способів усунення невизначених значень. Вважаючи класифікаційне правило наближеною функціональною залежністю із визначеним ступенем довіри A , використаємо для цього метод, аналогічний до відомого методу прогонки [2]: рівність значень атрибутів у лівій частині правила з ступенем довіри A означає і рівність значень атрибутів у правій частині.

Опишемо алгоритм застосування модифікованого методу прогонки.

Нехай у відношенні r підтримується наближена функціональна залежність $e(X_1, \dots, X_n \rightarrow A)$. Символ $\downarrow$ позначає визначене значення, а $\perp$ – його відсутність; t_i – кортеж відношення r (послідовність кортежів значення не має).

1. Якщо $\{t_1(X_1) \downarrow, \dots, t_1(X_n) \downarrow\} \wedge \{t_2(X_1) \downarrow, \dots, t_2(X_n) \downarrow\} \wedge \{t_1(X_1) \downarrow, \dots, t_1(X_n) \downarrow = t_2(X_1) \downarrow, \dots, t_2(X_n) \downarrow\} \wedge \{t_1(A) \downarrow\} \wedge \{t_2(A) = \perp\}$, то заміняємо кожне входження $\perp$ у r на $t_1(A)$.
2. Якщо $\{t_1(X_1) \downarrow, \dots, t_1(X_n) \downarrow\} \wedge \{ \text{в } t_2 \text{ } m \text{ з } n \text{ значень атрибутів} - \downarrow, n - m \text{ значень атрибутів} - \perp, m \leq n \} \wedge \{e \leq \frac{m}{n}\} \wedge \{ \text{за визначеними значеннями } t_1(X^m) \downarrow = t_2(X^m) \downarrow \} \wedge \{t_1(A) \downarrow\} \wedge \{t_2(A) = \perp\}$, то заміняємо кожне входження $\perp$ у r на $t_1(A)$.
3. Якщо $\{ \text{в } t_i \text{ } m_i \text{ з } n \text{ значень атрибутів} - \downarrow, m_i \leq n \} \wedge \{ \text{в } t_j \text{ } m_j \text{ з } n \text{ значень атрибутів} - \downarrow, m_j \leq n \} \wedge \{ \text{за визначеними значеннями } t_i(X^m) \downarrow = t_2(X^m) \downarrow \} \wedge \{ \text{за визначеними значеннями } t_j(X^m) \downarrow = t_2(X^m) \downarrow \} \wedge \{ \frac{m_i}{n} \leq \frac{m_j}{n} \} \wedge \{t_i(A) \downarrow\} \wedge \{t_j(A) \downarrow\} \wedge \{t_2(A) = \perp\}$, то заміняємо кожне входження $\perp$ у r на $t_j(A)$.

Наведемо приклад для демонстрації зменшення невизначеності у деякому відношенні, що описує обсяги продажу (невизначеність позначатимемо $\perp$).

Дата	Тип документа	Менеджер	Підрозділ	Клієнт	Продукт	К-ть	Ці-на	Сума
12.02.05	Про торгівлю	Іваненко	Склад 1	Самійленко	Дошка 0.5	10	2	20
12.02.05	Про торгівлю	Іваненко	Склад 1	Самійленко	Клей ПВА	2	1	2
12.02.05	Про торгівлю	Іваненко	⊥	Самійленко	Кромка АМ	2	2	4
12.02.05	Про торгівлю	Петренко	Склад 2	Литвиненко	Дошка 0.5	10	2	20
12.03.05	Про торгівлю	Петренко	Склад 2	Литвиненко	Дошка 0.5	10	2	20
12.04.05	Про торгівлю	Петренко	Склад 2	⊥	Дошка 0.5	10	2	20
12.05.05	Про торгівлю	Петренко	Склад 2	Литвиненко	Дошка 0.5	10	2	20

Експертами визначено такі класифікаційні правила:
 Тип документа, Менеджер, Продукт, Клієнт →Підрозділ.
 Тип документа, Менеджер, Підрозділ→Клієнт

Розглянемо, як здійснюється зменшення невизначеності.

У правій частині правил містяться атрибути *Підрозділ* та *Клієнт* з невідомими значеннями. Аналізуючи значення атрибутів, що містяться у лівій частині першого правила по перших двох кортежах, бачимо, що значення підрозділу необхідно встановити у *Склад1*, а значення атрибута *Клієнт* по кортежах 4,5,7 – у *Литвиненко*.

D. Структури даних для збереження класифікаційних правил

Класифікаційні правила (або нечіткі функціональні залежності) доцільно зберігати у окремому відношенні (словнику), можливий варіант схеми якого показано нижче.

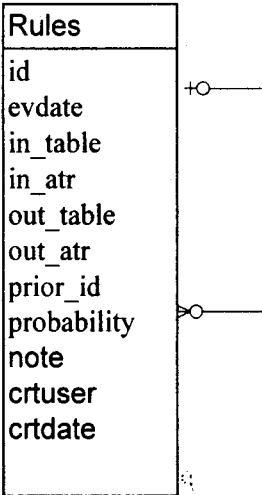


Схема відношення Rules

Схема відношення: ID – код, EVDATE – дата актуальності правила, IN_TABLE – назва таблиці-предка, IN_ATR – атрибут-предок, OUT_ATR атрибут-нащадок, OUT_TABLE – таблиця-нащадок, PRIOR_ID – зовнішній ключ таблиці Rules (для формування правил зі складеними частинами предків чи нащадків), PROBABILITY – довіра до правила.

Модифікований метод прогонки почергово перебирає усі правила з відношення Rules та застосовує його до кортежів відношень, вказаних у відповідному кортежі за правилом, що застосовується.

Висновки. Опрацювання невизначеності є ключовим моментом для багатьох методів видобування даних. Відомі сьогодні методи усунення невизначеностей опрацьовують лише відсутність, неповноту та нечіткість.

Наукова новизна. Запропоновано формальну класифікацію типів невизначеності, модель класу та класифікаційних правил як нечітких функціональних залежностей. Пропонуються методи визначення міри довіри до об’єктів класів.

Практична цінність. Наукові результати, отримані в цій статті, дають змогу вести подальші практичні дослідження по методах класифікації для усунення невизначеності.

1. Кравець Р.Б. Організація багатовимірного подання та аналізу інформації у реляційній базі даних // Вісник НУ "Львівська політехніка". 2003. № 489. 2. Netz A., Chaudhuri S. *Integration of Data Mining and Relational Databases. Proceedings of the 26th International Conference on Very Large Databases, Cairo, Egypt, 2000.* 3. G. Grsrfe, U. Fayyad. *On the Efficient Gathering of Sufficient Statistics for Classification from Large SQL Databases.* – 1998. – www.aaai.org. 4. Huhtala Y., Karkainen J. *Tane: An Efficient Algorithm for discovering Functional and Approximate Dependencies*// *The Computer Journal.* 1999. – Vol. 42. – № 2. 5. Дюбуа А., Прад А. Теория вероятностей. Приложение к представлению знаний в информатике. М., 1990. 6. Мейер Д. Теория реляционных баз данных. – М., 1987. 7. Шаховська Н.Б. Застосування апарату багатозначно. Логіки у системах баз даних // Вісник НУ "Львівська політехніка". 2001. № 438. 8. Шаховська Н.Б. Методи усунення невизначеностей у базах знань, побудованих на основі реляційного підходу // Вісник НУ "Львівська політехніка". – 2001. – № 438.

УДК 004.272.44

І.Д. Яковлєва, І.Д. Лісовенко

Чернівецький національний університет імені Юрія Федьковича
кафедра комп'ютерних систем та мереж

ПІДХІД ДО ПОБУДОВИ ПОТОКОВОГО ГРАФА АЛГОРИТМУ

© Яковлєва І.Д., Лісовенко І.Д., 2005

Розглянуто метод визначення готовності даних з використанням концепції машини потоків даних та запропоновано підхід до побудови потокового графа алгоритму.

In this paper the method of definition of data readiness based on the dataflow machine conception is considered and the approach to construction of algorithm flowgraph is offered.

Вступ. Для послідовних комп'ютерів основна модель виконання – відома модель фон Неймана, яка складається з послідовного процесу, що виконується в лінійному адресному просторі. Частка доступного паралелізму в ній порівняно мала. Принцип машини потоків даних (МПД), який подано у [1, 2], є одним з архітектурних шляхів досягнення високої продуктивності обчислень для максимального використання паралелізму.

Представлення алгоритмів у формі стандартного математичного запису не дає змоги достатньо повно оцінити можливості розпаралелювання та знайти способи їхньої ефективної реалізації. Як зазначено в [3], однією з форм представлення алгоритму, яка дає змогу оцінити його обчислювальні і структурні характеристики, є зображення алгоритму у вигляді функціонального графа.

Модель обчислювальної системи, що заснована на принципі потоку даних, автоматично використовує паралелізм, притаманний розв'язуванню задачам. Це значно полегшує організацію паралельної роботи, розподіл завдань при програмуванні і організацію обміну інформацією. В результаті продуктивність обчислювальної системи істотно зростає.

Мета цієї роботи – створити метод, який визначає готовність даних в послідовній програмі з використанням концепції МПД та побудувати потокову форму графа готовності даних алгоритму, яка описана в [5].

Описання алгоритму визначення готовності даних з використанням концепції МПД та побудова ярусно-паралельної форми алгоритму. В МПД команда оголошується готовою до