

ЛОГІКА КАУЗАЛЬНОГО ВИВЕДЕННЯ З ДАНИХ В УМОВАХ ПРИХОВАНИХ СПІЛЬНИХ ПРИЧИН

Анотація. Розглянуто проблеми виведення каузальних моделей з емпіричних даних і деякі механізми виникнення помилок. Показано, що відомі правила ідентифікації орієнтацій (спрямувань) статистичних зв'язків у каузальних моделях можуть втрачати адекватність, коли діють латентні конфаундери. Запропоновано корекції цих правил орієнтації, необхідні для їхнього застосування до моделей поза межами класу анцестральних моделей. Сформульовано необхідні припущення, які обґрунтовують виведення адекватних каузальних відношень з даних.

Ключові слова: каузальні моделі, d-сепарація, умовна незалежність, правила орієнтації ребер, конфаундер, колізор, ілюзорне ребро, припущення тестабільності залежності.

ВСТУП

Виявлення каузальних відношень та оцінювання каузального ефекту традиційно виконувалися на основі аналізу результатів рандомізованих експериментів. Натомість сучасні комп'ютерні методи глибокого аналізу великих даних обробляють дані, отримані як пасивні спостереження за середовищем (об'єктом). В умовах дефіциту предметних апіорних знань і відсутності темпоральної інформації задача виявлення каузальних відношень у даних концептуально важка. Цей напрямок досліджень інтенсивно розвивається впродовж останніх трьох десятиліть. Нині для розв'язання цієї задачі використовують переважно апарат каузальних мереж та марковських властивостей [1–6]. Для розпізнавання спрямованості (орієнтації) зв'язків моделі розроблено набір правил, які призначені для методів виведення моделі, оснований на незалежності [2, 7–9]. Існування прихованих спільних причин суттєво ускладнює розв'язання вказаних задач.

У цій статті показано, що деякі відомі правила орієнтації зв'язків моделі потребують уточнення і корекції. Для коректного застосування правил орієнтації ребер поза класом анцестральних моделей пропонуються оновлені версії правил. Уточнюються і пояснюються припущення, які забезпечують коректність та надійність виведення каузальних зв'язків.

1. КАУЗАЛЬНІ МОДЕЛІ ТА ЇХНІ ВЛАСТИВОСТІ

Емпірична каузальна модель адекватно відображає схему розгортання системи вимірних характеристик об'єкта. Найпоширенішим типом каузальних моделей є каузальні мережі, які специфікують структуровану систему спрямованих впливів між змінними. Кожній змінній відповідає вершина мережі, а кожному безпосередньому статистичному зв'язку — ребро. Одноорієнтоване ребро $X \rightarrow Y$ означає, що змінна (вершина) X є безпосередньою причиною для Y . Один кінець ребра $X \rightarrow Y$ (біля Y) названо «вістря», а другий (лівий) — «хвіст». Каузальна мережа — це пара (G, Θ) , де G — орграф, а Θ — параметри, асоційовані з G , які описують кількісний аспект моделі. Зазвичай

використовують моделі, де граф G орациклічний, тобто не містить циклів вигляду $X \rightarrow Y \rightarrow \dots \rightarrow X$. Якщо G орациклічний і утворений виключно з односторонніх ребер, то модель (G, Θ) належить класу оАОГ-моделей (на основі ординарних орграфів). Кількісний опис оАОГ-моделі складається з фрагментів у формі умовного розподілу $p(Y | \Phi(Y))$, де $\Phi(Y)$ — це сукупність «батьків» (безпосередніх причин) змінної Y . Зокрема, якщо маємо адитивний шум ε_Y , фрагмент моделі описується рівнянням вигляду $y = f(x, z, \dots) + \varepsilon_Y$.

Зазвичай деякі причини відсутні в заданому наборі даних, тобто приховані. Ситуація ускладнюється, коли прихована причина U впливає одночасно на дві або більше спостережуваних змінних. Це «нетривіальна» прихована причина. Далі будемо вважати, що нетривіальні приховані причини не мають вхідних ребер (вигляду $\rightarrow U$). Каузальна модель з нетривіальними прихованими причинами використовує біорієнтовані (двобічно-орієнтовані, або асоціативні) ребра вигляду $Q \leftrightarrow W$. Біорієнтоване ребро $Q \leftrightarrow W$ відображає вплив прихованої змінної одночасно (паралельно) на Q та W . Модель, що використовує і односторонні, і біорієнтовані ребра, називається мішаною (мАОГ-моделлю). Оскільки працювати із загальним випадком мАОГ-моделей складно, зазвичай обирають певний підклас цих моделей, визначений відповідними структурними обмеженнями.

Колізор (а collider) — це фрагмент вигляду $X * \rightarrow Y \leftarrow * Z$. Позначка $*$ в описі ребра $* \rightarrow$ є еквівалентним символом, який показує, що тип (вістря чи хвіст) цього кінця ребра є довільним або невідомим. Можливі чотири варіанти колізора: $X \rightarrow Y \leftarrow Z$; $X \leftrightarrow Y \leftarrow Z$; $X \rightarrow Y \leftrightarrow Z$; $X \leftrightarrow Y \leftrightarrow Z$. Колізор $X * \rightarrow Y \leftarrow * Z$ називатимемо шунтованим, якщо у графі є ребро між X та Z , інакше — нешунтованим. Шлях — це послідовність сусідніх ребер без повторення проміжних вершин. Ланцюг (безколізорний шлях) — це шлях, на якому немає жодного колізора. Якщо існує оршлях $X \rightarrow \dots \rightarrow Y$, то X є предком (причиною) для Y , а Y — нащадком для X . Цикл — це шлях, де перша й остання вершини тотожні. Циклон — це циклічний оршлях.

Нехай V — множина всіх спостережуваних змінних. Сумісний розподіл ймовірностей $p(V)$ можна факторизувати згідно з ланцюговим правилом Баєса. У розподілі $p(V)$ можуть виконуватися умовні незалежності, що дає змогу спростити формулу для $p(V)$. Умовну незалежність змінних X та Y за умови S позначатимемо $\text{Ind}(X; S; Y)$, а залежність — $\neg \text{Ind}(Z; S; W)$. Серед усіх фактів умовної незалежності важливими є ті, що імплікуються структурою моделі й інваріантні до її параметризації. Такі умовні незалежності є марковськими. У моделі (G, Θ) із класу оАОГ чинне правило: відсутність ребра породжує умовну незалежність. Для оАОГ-моделі розподіл $p(V)$ факторизується: $p(V) = \prod_{X \in V} p(X | \Phi(X))$. Усі марковські властивості моделі можна отримати з графу моделі G , застосовуючи критерій d-сепарації [1].

Означення 1 (d-сепарація, розгорнуте формулювання).

1. Шлях π в орграфі G називають d-перекритим (d-блокованим) за допомогою (кондиціонування) множини вершин S , якщо і тільки якщо:

а) на шляху π існує фрагмент $Z \rightarrow$ з тим, що $Z \in S$;

б) на шляху π лежить хоча б один колізор $* \rightarrow W \leftarrow *$, причому $W \notin S$ і немає жодної такої $T \in S$, що існує оршлях $W \rightarrow \dots \rightarrow T$.

2. Якщо всі шляхи між вершинами X та Y є d -перекритими (або немає жодного шляху між X та Y), то вершини X та Y є d -сепарованими.

3. Шлях, який не є d -перекритим, є d -відкритим.

4. Якщо за кондиціонування множини S у графі G існує принаймні один d -відкритий шлях між X та Y , то вершини X та Y є d -з'єднаними для заданої S .

Той факт, що множина вершин S d -сепарує вершини X та Y , позначатимемо предикатом $DS(X; S; Y)$. Тоді S є d -сепаратором для пари X, Y . Зрозуміло, що застережено $X, Y \notin S$. Коли d -сепаратор порожній, будемо писати $DS(X;; Y)$.

Визначення d -сепарації — це пп. 1 і 2. Випадок, коли критерій d -сепарації не задоволено, трактують пп. 3 і 4. Наслідки невиконання d -сепарації менш категоричні. Легко зрозуміти п. 1а критерію d -сепарації. Якщо вершини X та Y поєднані ланцюгом π , то щоб перекрити π , треба заблокувати вершину на π . Зокрема, еквівалентні такі два факти:

- вершини A та B є безумовно d -з'єднаними ($\neg DS(A;; B)$);
- між вершинами A та B існує принаймні один ланцюг.

В означенні 1 п. 1б можна прокоментувати так. Нехай шлях π має вигляд $X \rightarrow W \leftarrow Z$. За порожньої умови шлях π є d -перекритим. Але якщо як умову використати вершину W або її нащадка, шлях π стає d -відкритим. (Колізійний стик ребер розблоковується завдяки кондиціонуванню, «провокується» залежність між X та Z .)

Значення критерію d -сепарації розкривається через теорему [10], яка стверджує: у кожній каузальній мережі для всіх $X, Y \notin S$ виконується імплікація

$$DS(X; S; Y) \Rightarrow \text{Ind}(X; S; Y). \quad (1)$$

Було би зручно, щоб також виконувалась імплікація, обернена до (1), тобто якби кожна умовна незалежність, що виконується в моделі, була марковською. В дійсності така імплікація виконується регулярно, за винятком особливих випадків. Тобто має місце закономірність

$$\forall X, Y \notin S: \text{Ind}(X; S; Y) \Leftrightarrow DS(X; S; Y). \quad (2)$$

Якщо в (2) замінити послаблену імплікацію ($\Leftrightarrow$) строгою імплікацією, то отримаємо твердження, відоме як припущення каузальної «не-оманливості» (Causal faithfulness assumption) в його найжорсткішій формі [2, 3, 11]. Запишемо (2) в еквівалентній формі:

$$\forall X, Y \notin S: \neg DS(X; S; Y) \Leftrightarrow \neg \text{Ind}(X; S; Y), \quad (3)$$

що інтерпретується так: d -відкритий шлях (шляхи) створює залежність.

Кожна незалежність, не санкціонована структурою моделі, є нестійкою і руйнується (перетворюється на залежність) внаслідок навіть незначних змін значень параметрів. Відомі такі випадки порушення припущення каузальної «не-оманливості»: невиконання транзитивності залежності на ланцюзі; взаємна анігіляція паралельних впливів; «масковані» ребра. Невиконання транзитивності залежності означає, приміром, що маємо фрагмент моделі $X \rightarrow Y \rightarrow Z \rightarrow W$, де спостерігається $\text{Ind}(X;; W)$ чи навіть $\text{Ind}(X;; Z)$. Феномен маскованого ребра [1, 5, 12] полягає в тому, що в генеративній моделі є колізор $X \rightarrow Y \leftarrow Z$, причому виконується $DS(X;; Z)$ та $\neg \text{Ind}(Y;; \{X, Z\})$, і однак, маємо факт $\text{Ind}(X;; Y)$.

(Сумісний вплив факторів X та Z на залежну змінну Y явно проявляється, а маргінальний вплив одного фактора не проявляється.) Наведені випадки можливі в моделях з дискретними змінними. Натомість у моделях з неперервними змінними легко налаштувати параметри так, щоб впливи через шляхи $X \rightarrow Q \rightarrow Y$ та $X \rightarrow Z \rightarrow W \rightarrow Y$ взаємно анігілювалися, що приведе до $\text{Ind}(X;;Y)$.

Для практичних застосувань і, зокрема, для зручності виведення моделей з даних в умовах прихованих змінних було запропоновано спеціальний підклас класу мАОГ-моделей — анцестральні моделі [13]. Цей клас базується на анцестральних графах. Анцестральні (або предкові) графи утворюються з одноорієнтованих та біорієнтованих ребер з обмеженням: якщо вершини A та B поєднані оршляхом (зокрема, ребром), то забороняється ребро $A \leftrightarrow B$. Тобто забороняються не лише циклони, а й «майже циклони». Майже циклон утворюється з циклона вигляду $X \rightarrow \dots \rightarrow Y \rightarrow Z \rightarrow X$, якщо замінити будь-яке його ребро біорієнтованим, наприклад, $X \rightarrow \dots \rightarrow Y \rightarrow Z \leftrightarrow X$. По суті, в анцестральних графах забороняються латентні конфаундери (сплутувачі), тобто приховані причини, які впливають одночасно на причину A та її наслідок B . (Строго кажучи, в класі анцестральних графів також застосовуються ребра, які є неорієнтованими в принципі. Такі ребра виникають внаслідок селекції даних [13] і тут не розглядаються.)

У класі оАОГ-моделей діє таке правило: якщо задані вершини X та Y не поєднані ребром в G , то для пари X, Y існує d-сепаратор. Але ця властивість, на жаль, не поширюється на анцестральні графи. В анцестральних графах можливі випадки, коли вершини X та Y не поєднані ребром, проте, для пари X, Y не існує жодного d-сепаратора. Відсутність сепаратора для несуміжних вершин є проблемою для методів виведення моделі, основаних на незалежності. Між відповідними вершинами виникає, так би мовити, «ілюзорне» ребро. Зазначимо, що в анцестральних графах феномен ілюзорних ребер дещо обмежений: ілюзорне ребро не може виникати між вершинами, поєднаними оршляхом (між предком та нащадком). Натомість у моделях з латентними конфаундерами ілюзорне ребро між предком та нащадком можливе.

2. ПРИНЦИПИ ВИВЕДЕННЯ КАУЗАЛЬНИХ МОДЕЛЕЙ З ДАНИХ

Розглянемо методи та алгоритми виведення моделі, основані на незалежності. Еталонними алгоритмами цього підходу вважають PC, FCI та їхні модифікації [2–8]. Типовий алгоритм виводить модель у три етапи:

1) відтворює сукупність неорієнтованих ребер моделі (для кожної пари змінних шукається сепаратор);

2) визначає спрямування ребер (ребра орієнтуються згідно з правилами);

3) обчислює параметри моделі відповідно до виведеної структури.

Базовий принцип роботи алгоритму на етапі 1: якщо для пари X, Y знайдено сепаратор, то ребро між цими змінними вилучається, якщо не знайдено — ребро між X та Y встановлюється (утримується). Таке просте правило відображає фундаментальний принцип: серед моделей, узгоджених з даними, обрати найпростішу. Пошук сепараторів відбувається у порядку зростання рангу тестів, тобто зростання кардинальності умов тестів незалежності. Відомі також додаткові принципи тестування ребер.

Спрямування багатьох зв'язків моделі відтворити не вдається (вістря або хвіст ребра не розпізнаються). Це пояснюється марковською еквівалентністю

відповідних структур. Тому для опису моделі потрібні додаткові позначки для кінців ребер. Ребро вигляду $X \circ \circ Y$ відображає ситуацію, коли каузальний характер цього зв'язку зовсім не визначений (таку ситуацію маємо на виході етапу 1 алгоритму). Ребро $X \circ \circ Y$ будемо називати неорієнтованим (хоча, точніше, це ребро з невідомою орієнтацією). З огляду на можливе існування прихованих спільних причин знадобляться також ребра з частково визначеною орієнтацією вигляду $V \circ \rightarrow Z$. Таке ребро резервує три можливі варіанти: $V \rightarrow Z$, $V \leftrightarrow Z$ або обидва ці ребра одночасно. Пара ребер $V \rightarrow Z$, $V \leftrightarrow Z$ називається «арка». Далі будемо вважати, що арок немає. Ребро вигляду $V \circ \rightarrow Z$ назвемо можливо каузальним.

Розглянемо правила орієнтації ребер. Усі правила орієнтації ґрунтуються на принципі: обрати найбільш природне і лаконічне пояснення системи залежностей. Алгоритм намагається застосувати кожне правило, допоки є можливість. Правила орієнтації виконуються у встановленому порядку. Для класу оАОГ-моделей необхідним і достатнім є набір із чотирьох правил орієнтації [7] (їм відповідають наведені далі правила $\mathfrak{R}0$, $\mathfrak{R}1$, $\mathfrak{R}2$ та $\mathfrak{R}3$). Ще одне правило ($\mathfrak{R}4^{(0)}$) стає необхідним тільки, коли «ззовні» надходять додаткові каузальні орієнтації ребер (як апіорні знання). Для анцестральних моделей відомі правила були модифіковані і додатково розроблені нові («витончені») правила [8]. Далі аналізується підмножина правил, представлених у [8], оскільки конструкції з ребрами, що є по суті неорієнтованими, не розглядаються. (Збережено нумерацію правил.) Для правил орієнтації ребер передбачається, що виведено адекватні неорієнтовані ребра і що виконуються припущення «не-оманливості» (у відповідних версіях).

Правило $\mathfrak{R}0$ можна записати у такому вигляді: для кожної трійки X, Y, Z такої, що є пара ребер $X * \circ Y * \circ Z$ і немає ребра $X * \circ Z$, орієнтувати $X * \rightarrow Y \leftarrow * Z$, якщо й тільки якщо Y не входить до складу сепараторного набору для пари X, Z .

По суті, це правило розпізнає нешунтовані колізори, тому можна назвати його «колізорним». Наведене формулювання правила $\mathfrak{R}0$ включене в контекст опису алгоритму FCI. «Сепараторний набір для пари X, Z » розуміється як перший знайдений емпіричний сепаратор для пари X, Z . Взагалі, всі правила орієнтації описано в [8] як контекстні (вбудовані в контекст виконання алгоритму FCI). Сенс поняття сепараторного набору неявно залежить від тактики пошуку сепараторів та інших характеристик та параметрів методу. (В [8] вжито характеристику «розумний» алгоритм.) Для концептуальної ясності аргументації та зменшення контекст-залежності формулювань оберемо такий план викладу. В цьому і наступному розділі спробуємо виразити правила орієнтації у самодостатній формі, для чого ідеалізуємо проблему. А згодом повернемося до реалістичної проблемної ситуації.

Отже, нехай на вхід алгоритму виведення подано «теоретичний» (модельний) розподіл ймовірностей $p(\mathbf{V})$. Кожний знайдений факт умовної незалежності (згідно з (2)) майже завжди буде відповідати ізоморфному факту d-сепарації. Колізорне правило $\mathfrak{R}0$ можна записати у формі

$$(X * \circ Y * \circ Z) \& \exists \mathbf{S} (Y \notin \mathbf{S}) : \text{Ind}(X; \mathbf{S}; Z) \Rightarrow X * \rightarrow Y \leftarrow * Z. \quad (4)$$

(Зазначимо, що для $\text{Ind}(X; \mathbf{S}; Z)$ мається на увазі $X, Z \notin \mathbf{S}$.) Квантор існування $\exists \mathbf{S}$ використано в (4), щоб правило не було прив'язане до контексту кон-

кретного алгоритму. Але цей квантор існування надає формулюванню зайвої загальності. Це значно підсилить ризики помилок, коли від ідеалізованої постановки задачі перейдемо до виведення моделі з реалістичних вибірок даних. Для більшої конструктивності й стислості формулювань потрібно розумно звужити коло сепараторів, ігноруючи ті сепаратори, які алгоритм напевно не буде випробовувати. Для цього треба детальніше розглянути тактику пошуку сепараторних наборів під час виконання етапу 1 алгоритму виведення моделі.

Уявімо, що алгоритм отримує коректні результати тестів, тобто знайдені факти незалежності відповідають фактам d -сепарації у структурі генеративної моделі. Наведемо й зіставимо три принципи пошуку сепараторів. «Наївний» принцип формує «гіпотетичні» сепаратори (сепаратори-кандидати) для пари X, Y із змінних, залежних від X (відповідно Y), не пропускаючи жодних комбінацій. Найбільш відомий («базовий») принцип (який широко застосовується, починаючи з алгоритму РС) полягає в тому, що пробні (гіпотетичні) сепаратори для пари X, Y формуються із змінних, можливо, суміжних з X (відповідно Y). Натомість «систематичний» принцип формує пробні сепаратори для пари X, Y згідно з необхідними вимогами до локально-мінімальних сепараторів та до їхніх членів (елементів). Ці вимоги сформульовано в [14–17]. Алгоритм виведення, побудований на апараті локально-мінімальної сепарації, називатимемо «систематичним». (Можна поєднати базовий та систематичний принципи в одному алгоритмі.)

Якщо для пари вершин X, Y існують сепаратори (сепаратор), то серед них є один або кілька мінімальних сепараторів (однакової кардинальності) і, можливо, ще кілька локально-мінімальних сепараторів різної кардинальності. (Кожний мінімальний сепаратор тривіально є локально-мінімальним.) Алгоритм, побудований за наївним принципом, завжди відшукує один з мінімальних сепараторів (якщо сепаратор для цієї пари існує). Систематичний алгоритм також завжди знаходить один з мінімальних сепараторів, але в процесі пошуку зазвичай виконує значно меншу кількість тестів. У разі, коли сепаратора не існує, систематичний принцип забезпечує ще більшу економію у тестах. Базовий принцип пошуку сепараторів також забезпечує значне зменшення кількості тестів (часто перевершуючи систематичний принцип). Проте, базовий принцип гарантує відшукання тільки локально-мінімального сепаратора, але іноді не знаходить мінімального сепаратора [14, 17]. Більше того, в анцестральних моделях (і взагалі в мішаних каузальних моделях) алгоритм, який дотримується тільки базового принципу, іноді не може відшукати жодного сепаратора, навіть коли потрібний сепаратор дійсно існує. Тож, щоб уникнути помилок, алгоритм FCI виконує два додаткових етапи. Коли далі буде вживатися поняття мінімального сепаратора, алгоритмом за умовчанням є систематичний. (Втім, часто алгоритм з базовим принципом відшукує ті самі сепаратори.)

Уявімо, що існує алгоритм A , який шукає всі локально-мінімальні сепаратори для заданої пари X, Y . Алгоритм A використовує правила локально-мінімальної сепарації для формування гіпотетичних сепараторів (сепараторів-кандидатів). Нехай $K(X, Y | G)$ — множина гіпотетичних сепараторів, сформованих згідно з цими правилами. Множина $K(X, Y | G)$ залежить від набору правил локально-мінімальної сепарації. Будемо вважати, що алгоритм має набір правил, достатній, щоб забезпечити значну економію тестів.

У цілому, алгоритм $\mathbf{A}$ випробує множину гіпотетичних сепараторів $\mathbf{K}(X, Y | G) \cup \{Null\}$, де $\{Null\}$ — порожня множина вершин (вона відповідає тесту безумовної незалежності).

Тепер повернемося до систематичних алгоритмів, які шукають тільки один сепаратор для кожної пари X, Y . Якщо в структурі G для пари вершин X, Y не існує жодного сепаратора, то для X, Y систематичний алгоритм випробує множину гіпотетичних сепараторів $\mathbf{K}(X, Y | G) \cup \{Null\}$. Якщо для пари вершин X, Y існують сепаратори (сепаратор), то систематичний алгоритм відшукає один із мінімальних сепараторів $dSep_{\min}(X, Y)$. Якщо мінімальний сепаратор не порожній, то в процесі пошуку алгоритм випробує множину сепараторів-кандидатів $\mathbf{\Lambda}(X, Y | G)$, причому потужність кожного набору $S_i \in \mathbf{\Lambda}(X, Y | G)$ обмежена: $|S_i| \leq |dSep_{\min}(X, Y)|$. Ясно, що $\mathbf{\Lambda}(X, Y | G) \subseteq \mathbf{K}(X, Y | G)$. Оскільки в цьому розділі аналізується виведення моделі з теоретичного розподілу $p(\mathbf{V})$, усе викладене раніше про d-сепарацію можна транслювати на факти умовної незалежності, знайдені алгоритмом (із застереженням особливих випадків). Зокрема, набору вершин $dSep_{\min}(X, Y)$ відповідає набір змінних $Sep_{\min}(X, Y) = S$, який забезпечує виконання $Ind(X; S; Y)$.

Повернемося до формулювань правила $\mathfrak{R0}$ та інших правил орієнтації. У контексті систематичного алгоритму можна замінити занадто теоретичне формулювання (4) правила $\mathfrak{R0}$ прагматичнішим. Отримуємо таку форму для $\mathfrak{R0}$:

$$X * \multimap Y * \multimap Z \ \& \ (X * \multimap Z \text{ is absent}) \ \& \ Y \notin Sep_{\min}(X, Z) \Rightarrow X * \multimap Y \leftarrow * Z. \quad (5)$$

Далі, коли у формулі фігурує терм $Y \notin Sep_{\min}(X, Z)$, за умовчанням вважається, що $Sep_{\min}(X, Z)$ існує.

Після повного застосування правила $\mathfrak{R0}$ (якщо воно спрацювало принаймні один раз) алгоритм виведення переходить до наступного правила $\mathfrak{R1}$ орієнтації ребер.

Правило $\mathfrak{R1}$ можна записати у двох формах (відповідно до стилю формулювань (4) та (5)):

$$\begin{aligned} X * \multimap Y \circ \multimap Q \ \& \ \exists S(Y \in S) : Ind(X; S; Q) \Rightarrow X * \multimap Y \rightarrow Q; \\ X * \multimap Y \circ \multimap Q \ \& \ Y \in Sep_{\min}(X, Q) \Rightarrow X * \multimap Y \rightarrow Q. \end{aligned} \quad (6)$$

Правило $\mathfrak{R2}$ має два варіанти умови:

$$X \rightarrow Y * \multimap W \ \& \ X * \multimap W \Rightarrow X * \multimap W; \quad (7)$$

$$X * \multimap Y \rightarrow W \ \& \ X * \multimap W \Rightarrow X * \multimap W. \quad (8)$$

Правило $\mathfrak{R3}$ можна записати так:

$$X * \multimap W \leftarrow * Z \ \& \ X * \multimap Y \circ \multimap Z \ \& \ Y * \multimap W \ \& \ \exists S(Y \in S) : Ind(X; S; Z) \Rightarrow Y * \multimap W.$$

Прагматичне формулювання правила $\mathfrak{R3}$ має вигляд

$$X * \multimap W \leftarrow * Z \ \& \ X * \multimap Y \circ \multimap Z \ \& \ Y * \multimap W \ \& \ Y \in Sep_{\min}(X, Z) \Rightarrow Y * \multimap W. \quad (9)$$

Правило $\mathfrak{R4}$ суттєво відрізняється від правила $\mathfrak{R4}^{(0)}$, призначеного для класу оАОГ-моделей. У формулюванні $\mathfrak{R4}$ використано поняття дискримінаційного шляху (discriminating path) [6, 8, 13]. Для спрощення не будемо використовувати це поняття. Розглянемо найпростішу структуру, яка задо-

вольняє умови правила $\mathcal{R}4$, і для неї сформулюємо локальний (спрощений) варіант правила, позначивши його $\mathcal{R}4_{(1)}$. Спробуємо описати це правило більш формально, водночас дотримуючись змісту оригінального опису правила $\mathcal{R}4$ [8]. Тоді отримаємо таке формулювання правила $\mathcal{R}4_{(1)}$:

Let $A \ast \rightarrow B \rightarrow C \ \& \ B \leftarrow \ast D \circ \ast C \ \& \ (A \text{ isn't adjacent to } C)$. Then

$$D \in \text{Sep}(A, C) \Rightarrow D \rightarrow C;$$

$$D \notin \text{Sep}(A, C) \Rightarrow B \leftrightarrow D \leftrightarrow C.$$

Тут $\text{Sep}(A, C)$ — це той сепараторний набір змінних для пари A, C , який був знайдений алгоритмом виведення. Запишемо правило $\mathcal{R}4_{(1)}$ у запропонованій раніше формі окремо для двох випадків:

$$A \ast \rightarrow B \rightarrow C \ \& \ B \leftarrow \ast D \circ \ast C \ \& \ \exists S: \text{Ind}(A; S; C), \ D \in S \Rightarrow D \rightarrow C; \quad (10)$$

$$A \ast \rightarrow B \rightarrow C \ \& \ B \leftarrow \ast D \circ \ast C \ \& \ \exists S: \text{Ind}(A; S; C), \ D \notin S \Rightarrow B \leftrightarrow D \leftrightarrow C.$$

Прагматичне формулювання правила $\mathcal{R}4_{(1)}$ має вигляд

$$A \ast \rightarrow B \rightarrow C \ \& \ B \leftarrow \ast D \circ \ast C \ \& \ D \in \text{Sep}_{\min}(A, C) \Rightarrow D \rightarrow C; \quad (11)$$

$$A \ast \rightarrow B \rightarrow C \ \& \ B \leftarrow \ast D \circ \ast C \ \& \ D \notin \text{Sep}_{\min}(A, C) \Rightarrow B \leftrightarrow D \leftrightarrow C. \quad (12)$$

Наступна за складністю структура, для якої призначене правило $\mathcal{R}4$, містить два колізори на шляху між вершинами A та D . Відповідні варіанти правила $\mathcal{R}4_{(2)}$ описуються так:

$$A \ast \rightarrow B \rightarrow C \ \& \ B \leftrightarrow T \rightarrow C \ \& \ T \leftarrow \ast D \circ \ast C \ \& \ D \in \text{Sep}_{\min}(A, C) \Rightarrow D \rightarrow C; \quad (13)$$

$$A \ast \rightarrow B \rightarrow C \ \& \ B \leftrightarrow T \rightarrow C \ \& \ T \leftarrow \ast D \circ \ast C \ \& \ D \notin \text{Sep}_{\min}(A, C) \Rightarrow T \leftrightarrow D \leftrightarrow C. \quad (14)$$

3. КОРЕКЦІЯ ПРАВИЛ ОРІЄНТАЦІЙ РЕБЕР ДЛЯ ЗАСТОСУВАННЯ ПОЗА МЕЖАМИ КЛАСУ АНЦЕСТРАЛЬНИХ МОДЕЛЕЙ

З практичної точки зору нереалістично сподіватися, що невідома генеративна модель належить класу анцестральних. Крім того, правила орієнтації ребер, описані в [8] і втілені в алгоритмі FCI, іноді наївно трактуються як універсальні правила для моделей з прихованими змінними, що може призвести до неадекватних висновків. Отже, є потреба розглянути можливість застосування вказаних правил орієнтації ребер поза межами класу анцестральних моделей і зробити необхідні корекції для такого застосування.

Окреслимо певний клас структур моделей. Структури цього класу утворюються з одноорієнтованих та біорієнтованих ребер і не мають циклонів, але майже циклони дозволені. Отже, дозволені латентні конфаундери. Встановимо обмеження для обраного класу — кожна нетривіальна прихована причина U впливає одночасно тільки на якісь дві спостережувані змінні: A та B , причому в генеративній моделі відсутнє ребро $A \rightarrow B$. Назвемо цей клас моделями без гіперребер та арок. Указане структурне обмеження не є конче необхідним для поставленої задачі, а лише надає конкретики. (Для цього класу працюють відомі методи розпізнавання ілюзорних ребер.)

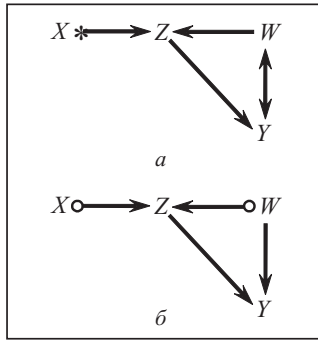


Рис. 1. Ілюстрація застосування правила $\mathcal{R}4_{(1)}$: генеративна модель (а); виведена модель (б)

Правила $\mathcal{R}0$ та $\mathcal{R}1$ не потребують змін для вказаного класу моделей.

Правила $\mathcal{R}2$, $\mathcal{R}3$ та $\mathcal{R}4$ обґрунтовані вимогою, щоб структура моделі була анцестральною [8, 9, 18, 19]. Дійсно, розглянемо, наприклад, правило $\mathcal{R}2$. Якщо виконуються умови правила (7), то заперечувати висновок $X * \rightarrow W$ цього правила означає стверджувати, що можлива орієнтація ребра $X \leftarrow W$. Але в разі орієнтації $X \leftarrow W$ (в залежності від уточнення ребра $Y * \rightarrow W$) виникне або цикл, або майже цикл $X \rightarrow Y \leftrightarrow W \rightarrow X$.

Припустимо, що для генеративної моделі з вказаного класу алгоритм встановив (на етапі 1) адекватний набір ребер. Оскільки генеративна модель може містити майже цикли, умови правил орієнтації мають бути жорсткішими. Після корекції маємо такі правила.

Замість (7) та (8) тепер буде правило $\mathcal{R}2^{(N)}$

$$X \rightarrow Y \rightarrow W \& X * \circ W \Rightarrow X * \rightarrow W \quad (15)$$

(верхній індекс N правила має асоціюватися зі словами «нове» або «неанцестральні»).

Правило $\mathcal{R}3^{(N)}$ має вигляд:

$$X \rightarrow W \leftarrow Z \& X * \circ Y \circ * Z \& Y * \circ W \& Y \in \text{Sep}_{\min}(X, Z) \Rightarrow Y * \rightarrow W. \quad (16)$$

Детальніше розглянемо застосування правила $\mathcal{R}4_{(1)}$ поза класом анцестральних структур. (Подальша логіка поширюється й на загальну версію правила $\mathcal{R}4$.) Нехай генеративна модель має структуру, показану на рис. 1, а. Марковськими властивостями цієї моделі є $\text{Ind}(X;; W)$ та $\text{Ind}(X; \{Z, W\}; Y)$. Застосуємо алгоритм з оригінальними (наведеними у розд. 2) правилами орієнтації. Правило $\mathcal{R}0$ виведе орієнтацію $X \circ \rightarrow Z \leftarrow \circ W$. Правило $\mathcal{R}1$ виконає орієнтацію $X \circ \rightarrow Z \rightarrow Y$. Правило $\mathcal{R}2$ орієнтує ребро $W \circ \rightarrow Y$. У цій ситуації задовольняються умови правила $\mathcal{R}4_{(1)}$, а саме випадку (11) правила. Тому буде виведено $W \rightarrow Y$. Бачимо, що застосування правила $\mathcal{R}4_{(1)}$ породжує неадекватне ребро (замість біорієнтованого виводиться каузальне ребро). Виведена модель показана на рис. 1, б. (Детальніше про застосування правила $\mathcal{R}4_{(1)}$ див. у [20].) Отже, випадок (11) правила $\mathcal{R}4_{(1)}$ потребує корекції.

Корегована версія $\mathcal{R}4_{(1)}^{(N)}$ правила має такий вигляд:

$$A * \rightarrow B \rightarrow C \& B \leftarrow * D \circ * C \& D \in \text{Sep}_{\min}(A, C) \Rightarrow B \leftarrow * \underline{D} \circ \rightarrow C; \quad (17)$$

$$A * \rightarrow B \rightarrow C \& B \leftarrow * D \circ * C \& D \notin \text{Sep}_{\min}(A, C) \Rightarrow B \leftrightarrow D \leftrightarrow C. \quad (18)$$

У правилі (17) після знаку імплікації замість позначки хвоста використано позначку невизначеного кінця ребра $D \circ \rightarrow C$. Конструкція $\leftarrow * \underline{D} \circ \rightarrow$ вказує на неколізорний стик ребер.

Взявши за основу традиційний алгоритм виведення, замінимо $\mathcal{R}2$ правилом $\mathcal{R}2^{(N)}$, а правило (11) правилом (17). Тоді для наведеного прикладу

(рис. 1, *a*) правило $\mathcal{R}2^{(N)}$ не діє, а правило $\mathcal{R}4_{(1)}^{(N)}$ виведе ребро $W \circ \rightarrow Y$, що не суперечить генеративній моделі. Невизначений кінець цього ребра маскує той факт, що модель виходить за межі класу анцестральних.

Для того самого набору зв'язків (див. рис. 1, *a*) спробуємо побудувати сценарій виведення, який приведе до точнішої орієнтації ребер, так, щоб результатом була неанцестральна модель. Створимо умови для розпізнавання біорієнтованого ребра між W та Y . Для цього розширимо модель і утворимо нешунтований колізор, який би дав змогу правилу $\mathcal{R}0$ встановити друге вістря ребра $W \leftrightarrow Y$. Включимо в генеративну модель додатковий фрагмент із трьома змінними: Q , T та V . У результаті матимемо генеративну модель, показану на рис. 2, *a*. (Формально тепер маємо іншу задачу виведення. Але по суті модель, показана на рис. 1, *a*, може бути отримана з моделі, показаної на рис. 2, *a*, якщо приховати змінні Q , T , V .) З'ясовується, що не вдається вивести модель з вочевидь неанцестральною структурою. Дійсно, сепарувати змінні Q та Y неможливо. Ланцюг $Y \leftarrow Z \leftarrow W \leftarrow * Q$ забезпечує залежність між Q та Y . Щоб перекрити цей ланцюг, треба кондиціонувати змінну (вершину) W або Z , але це відкриє шлях через колізор генеративної моделі $Y \leftrightarrow W \leftarrow * Q$. Отже, для пари Q , Y не існує сепаратора і буде виведене ілюзорне ребро $Y \circ \rightarrow Q$. Далі будуть виведені орієнтації $X \circ \rightarrow Z \leftarrow \circ W$ та $X \circ \rightarrow Z \rightarrow Y$. У «доданій» частині моделі правило $\mathcal{R}0$ встановить орієнтації $T \circ \rightarrow Q \leftarrow \circ V$. Унаслідок наявності ребра $Y \circ \rightarrow Q$ у виведеній моделі колізор $Y \leftrightarrow W \leftarrow * Q$ стає немов би «шунтованим». Тому правило $\mathcal{R}0$ для нього не діє. Кінець ребра $W \circ \rightarrow Y$ залишається невизначеним. Потім правило $\mathcal{R}1$ виведе каузальні ребра $W \leftarrow Q$ та $Y \leftarrow Q$, а далі (на відміну від попереднього випадку) правило $\mathcal{R}1$ виведе каузальне ребро $Z \leftarrow W$. Тепер спрацює правило $\mathcal{R}2^{(N)}$ і встановить орієнтацію $W \circ \rightarrow Y$. На цьому процес орієнтацій ребер завершується. Отже, не вдалося вивести біорієнтоване ребро $W \leftrightarrow Y$. Завдяки ілюзорному ребру факт неанцестральності моделі залишається замаскованим. Але якщо в алгоритмі замість (17) застосовується оригінальне правило $\mathcal{R}4_{(1)}$, наприклад (11), то це правило орієнтує ребро як $W \rightarrow Y$, що є неадекватним.

Оскільки ілюзорне ребро у виведеній моделі отримало статус каузального, воно може породжувати помилкові каузальні висновки. Візуальний аналіз виведеної моделі (див. рис. 2, *б*) дає аналітику підстави вважати, що змінна Q має ще окремий «паралельний» вплив на змінну Y крізь ребро $Q \rightarrow Y$ (і не виключено, що існує ще й третій оршлях впливу на Y в обхід Z , а саме по можливому оршляху $Q \rightarrow W \rightarrow Y$). Але це не так. Припустимо, генеративна модель відома. Тоді ясно, що керування змінною Z справить певний ефект на Y і цей ефект не зміниться від одночасного керування змінною Q . Формально це можна виразити в апараті числення каузального ефекту (*do-calculus*) [1]. Якщо розглядати

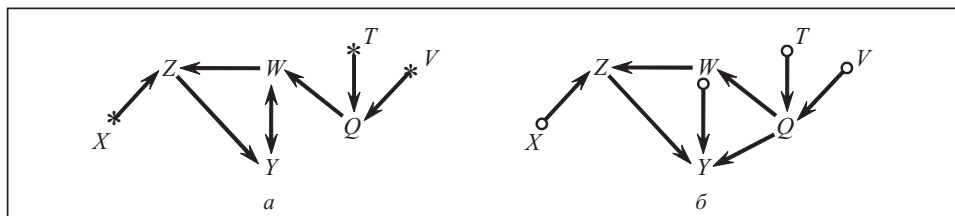


Рис. 2. Приклад виведення ілюзорного ребра $Q \rightarrow Y$: генеративна модель (*a*); виведена модель (*б*)

виведену модель (див. рис. 2, б), то завдяки каузальному ілюзорному ребру $Q \rightarrow Y$ складається враження, що має бути $p(Y | do(Z), do(Q)) \neq p(Y | do(Z))$. Але генеративна модель показує, що ці два ефекти ідентичні.

Наведені аргументи й висновки отримано для ідеалізованої проблеми, тобто ґрунтуються на жорсткому припущенні каузальної «не-оманливості». Як показано в [20], для подібної генеративної моделі в разі емпіричного виведення моделі з даних може бути виведена неанцестральна модель. Існує кілька сценаріїв, що приводять до вилучення ілюзорного ребра $Y \circ\!\!\!\circ Q$. Це може статися завдяки порушенню припущення каузальної «не-оманливості» (див. розд. 5). Вилучити ілюзорне ребро можна також за допомогою стримування Верма [1, 13, 21–24]. І насамкінець, вилучити ілюзорне ребро можна на підставі специфічного предметного знання. Наприклад, нехай маємо апіорні предметні знання, що в генеративній моделі змінна Q має тільки один механізм впливу на змінну Y (тобто не існує «паралельних» шляхів впливу). Нехай з даних виведено модель, відображену на рис. 2, б. Тоді, оскільки вилучення ребра $Q \rightarrow W$, $W \rightarrow Z$ чи $Z \rightarrow Y$ призводить до неадекватності моделі, то вилучити треба ребро $Q \rightarrow Y$.

Зазначимо, що серед сформульованих у цьому розділі правил єдиним, що виводить каузальний зв'язок, є правило $\mathfrak{R}1$.

4. ПРАГМАТИКА ВИВЕДЕННЯ МОДЕЛІ З ДАНИХ. РИЗИКИ ПОМИЛОК

У реальних задачах на вхід алгоритму виведення моделі подається вибірка даних. Вибірковий розподіл ймовірностей $p^{(\approx)}(\mathbf{V})$ відрізняється від «теоретичного» розподілу $p(\mathbf{V})$. Внаслідок вибіркового ухилення виникають емпіричні залежності між теоретично незалежними змінними. Для того щоб виведена з даних модель мала якнайменше помилкових (зайвих) ребер, потрібно явно слабку статистичну залежність трактувати як незалежність. Водночас вибіркоче ухилення може послаблювати статистичні залежності на позиціях деяких автентичних ребер генеративної моделі. (Навіть за відсутності шуму деякі залежності є слабкими завдяки відповідним значенням параметрів.) У викладеному раніше (див. коментарі до (3)) вже згадувалися механізми, які породжують немарковські незалежності в теоретичному розподілі $p(\mathbf{V})$. Такі події характеризуються нульовою мірою за Лебегом [11]. Натомість слабкі залежності є типовими (хоча й породжуються аналогічними механізмами). Досить часто тест показує незалежність між теоретично залежними змінними. Щоб не вилучати автентичні ребра моделі, треба утриматися від завищених вимог до сили залежностей. Отже, відповідність між емпіричними залежностями та структурою моделі розмивається і спотворюється. У процесі виведення моделі об'єктивно виникає ризик помилок. Прагматичне розв'язання полягає у відшуванні балансу помилок двох типів (втрата ребра — зайве ребро). Але для малих вибірок даних такий баланс не розв'язує проблеми, бо виникає багато помилок обох типів одночасно. Тоді треба взагалі відмовитися від спроби вивести каузальну модель. Далі будемо вважати, що маємо достатньо великі вибірки даних.

Зниження ризику помилок досягається сукупністю засобів, які включають: «розумний» алгоритм пошуку сепараторів, оптимальний вибір техніки обчислення тестової статистики, оптимальний вибір параметрів тестів тощо. Баланс помилок двох типів досягається за рахунок вибору рівня «альфа» для тесту. Будемо вважати, що викладені раніше вимоги враховано і, отже, вико-

ристовуються «акуратні» тести незалежності. Позитивний результат акуратного тесту незалежності (коли тест не відхиляє гіпотези незалежності) позначатимемо $\text{Ind}_{\text{Emp}}(X; S; Y)$. Тобто $\text{Ind}_{\text{Emp}}(*; *; *)$ означає емпіричну незалежність. Відповідно $\neg \text{Ind}_{\text{Emp}}(*; *; *)$ — це негативний результат тесту, емпірична залежність. Емпіричну версію множини сепараторів-кандидатів $\mathbf{K}(X, Y | G)$ позначимо $\mathbf{K}_{\text{Emp}}(X, Y)$. Аналогічно вводиться $\Lambda_{\text{Emp}}(X, Y)$.

Перша необхідна вимога до акуратних тестів — не помилятися у випадках теоретичної незалежності. Цю вимогу можна записати формально як імплікацію

$$\forall X, Y \notin S: [\text{Ind}(X; S; Y) \Rightarrow \text{Ind}_{\text{Emp}}(X; S; Y)].$$

Утім, на практиці можна послабити вимогу. В цій формулі квантор загальності є необов'язковим. Немає нагальної потреби розглядати усі можливі факти незалежності та розглядати всі сепаратори S . Для відповідних алгоритмів достатньо розглядати лише мінімальні або локально-мінімальні сепаратори. Тоді вимога, що забезпечує алгоритму коректне вилучення ребер, зводиться до імплікації

$$\forall X, Y: [S = \text{Sep}_{\min}(X, Y) \Rightarrow \text{Ind}_{\text{Emp}}(X; S; Y)].$$

Водночас потрібно, щоб акуратний тест незалежності давав негативний результат у всіх «критичних» випадках теоретичної залежності. Ці критичні випадки конкретизуються прагматичними версіями припущення каузальної «не-оманливості», на які спирається обґрунтування алгоритму [2, 5, 17, 25–28].

Якщо жорстку форму припущення каузальної «не-оманливості» буквально перевести в емпіричну форму, то припущення буде виражатися як імплікація $\forall X, Y \notin S: \neg \text{DS}(X; S; Y) \Rightarrow \neg \text{Ind}_{\text{Emp}}(X; S; Y)$. Таку вимогу можна назвати припущенням повної прагматичної каузальної «не-оманливості». Це припущення нереалістичне і його можна послабити. Для обґрунтування коректності методів виведення моделі буде достатньо й менш жорстких припущень. Запропонуємо необхідні й обережні версії прагматичних припущень-вимог. Будемо називати ці прагматичні версії припущеннями тестабільності.

Буває, що тест з умовою S фіксує емпіричну незалежність несуміжних змінних X та Y , попри те, що деякі (довгі й слабкі) шляхи між X та Y залишаються відкритими з цим S . Тоді маємо $\Lambda_{\text{Emp}}(X, Y) \subset \Lambda(X, Y | G)$. У деяких інших ситуаціях (коли між X та Y є ілюзорне ребро) графового сепаратора для X, Y не існує, але існує «нелегітимний», втім, толерантний емпіричний сепаратор.

Для обґрунтування етапу 1 виведення моделі потрібно таке «первинне» припущення: якщо в генеративній моделі існує ребро $X \leftrightarrow Y$, то всі тести умовної незалежності для пари X, Y мусять дати негативний результат $\neg \text{Ind}_{\text{Emp}}(X; S; Y)$. Це припущення можна формально записати як імплікацію

$$X \leftrightarrow Y \Rightarrow \forall S(X, Y \notin S): \neg \text{Ind}_{\text{Emp}}(X; S; Y). \quad (19)$$

Таке припущення (з назвою Adjacency-Faithfulness) запропоновано в [25]. Утім, квантор загальності в (19) для типових алгоритмів не є обов'язковим і встановлює завищені вимоги до реберної залежності. (У розподілі $p^{(\approx)}(\mathbf{V})$ можна знайти багато аномальних паттернів, які не є критичними.) Для систематичних алгоритмів достатньо, щоб тест незалежності давав негативний результат для наборів змінних S , сформованих згідно з правилами пошуку локально-мінімальних сепараторів для пари X, Y [16, 17]. Отже, отримуємо

прагматичне припущення тестабільності реберної залежності (ПТРЗ):

$$X \circ - * Y \Rightarrow \forall S \in [\mathbf{K}_{\text{Emp}}(X, Y) \cup \{\text{Null}\}]: \neg \text{Ind}_{\text{Emp}}(X; S; Y). \quad (20)$$

Виконання (20) є другою необхідною вимогою до акуратних тестів. Для обґрунтування правил орієнтації ребер необхідні інші версії прагматичних припущень-вимог; вони будуть викладені у розд. 6. Строго кажучи, для того щоб сформульовані раніше правила орієнтації застосовувати практично, їх треба було би переписати, замінивши конструкт $\text{Sep}_{\min}(*, *)$, визначений для теоретичного розподілу $p(\mathbf{V})$, емпіричним аналогом. (Це стосується правил (5), (6), (9), (11)–(14), (16)–(18). Правило (15) не містить цього конструкта.)

5. ДЕМОНСТРАЦІЯ РИЗИКІВ ТА АНОМАЛІЙ ОРІЄНТАЦІЇ РЕБЕР У ВИВЕДЕННІ МОДЕЛІ З ЕМПІРИЧНИХ ДАНИХ

Покажемо, що виведення моделі з даних стандартними алгоритмами може приводити до альтернативних результатів. Зокрема, для генеративної моделі за межами класу анцестральних графів знайдено сценарії, коли коректний алгоритм (наприклад, FCI) виводить адекватну модель (неанцестральність якої репрезентовано явно). Було проведено низку чисельних експериментів (див. [20]). Використовувалась лінійна модель зі структурою, показаною на рис. 3, а. Якщо результати всіх виконаних тестів незалежності узгоджуються з фактами d-сепарації, то буде виведено ілюзорне ребро $Y \circ - \circ Q$, і тоді правило (17) буде орієнтувати ребро $W \circ - \rightarrow Y$. Невизначений кінець цього ребра маскує той факт, що генеративна модель не є анцестральною. Натомість у [20] знайдено доволі реалістичні сценарії, коли ілюзорне ребро вилучається під час виведення. За одним сценарієм гіпотеза незалежності Q та Y утримується завдяки тому, що безумовна кореляція між змінними Q та Y слабка, за іншим — завдяки тому, що за умови на Z частинна кореляція між змінними Q та Y слабка. Внаслідок вилучення ребра $Y \circ - \circ Q$ правило $\mathcal{R}0$ встановить орієнтації $Y \circ - \rightarrow W \leftarrow \circ Q$. Також $\mathcal{R}0$ орієнтує $X \circ - \rightarrow Z \leftarrow \circ W$. Правило $\mathcal{R}1$ орієнтує $Z \rightarrow Y$ та $Z \leftarrow W$. У цій ситуації буде застосоване правило $\mathcal{R}2^{(N)}$, що дасть $W \leftrightarrow Y$. Отже, у виведеній моделі з'явиться майже цикл $W \rightarrow Z \rightarrow Y \leftrightarrow W$. На цьому виведення завершиться. Результат показано на рис. 3, б. (Правило (17) не працює для фрагмента зі змінними X, Z, W, Y , оскільки умови правила вимагають $W \circ - * Y$. Для іншого фрагмента, який складається з трикутника ребер та ребра $W \leftarrow \circ Q$, правило (18) не працює, тому що не задоволено умову $Y \circ - * Z$.) Отже, виведено адекватну неанцестральну модель.

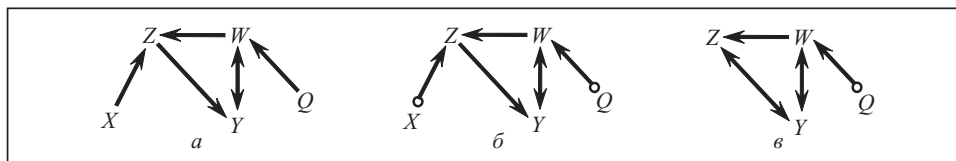


Рис. 3. Ілюстрація виведення моделі з даних, коли ілюзорне ребро вилучається: генеративна модель (а); виведена модель (б); виведена модель, коли змінна X прихована (в)

Якщо змінити постановку задачі, а саме зробити змінну X прихованою, а все інше (кореляції, алгоритм, поріг значущості тестів) зберегти незмінним, то залежність між Q та Y залишиться слабкою. Тому також буде вилучене ілю-

зорне ребро. Але тепер після правил $\mathcal{R}0$ та $\mathcal{R}1$ спрацює правило (18), яке встановить орієнтації $Z \leftrightarrow Y \leftrightarrow W$. Виведена модель показана на рис. 3, в. Орієнтація ребра $Z \leftrightarrow Y$ є помилковою.

Отже, порушення припущення каузальної «не-оманливості» (у формі слабкості залежності, утвореної трьома ребрами) може призводити до аномальних результатів. В одних випадках отримуємо адекватну модель поза межами анцестральних графів, у інших — біорієнтоване ребро на позиції каузального зв'язку. Для обґрунтування коректності правил орієнтацій потрібно уточнювати припущення.

6. УМОВИ НАДІЙНОСТІ ОРІЄНТАЦІЙ РЕБЕР

У розд. 4 сформульовано прагматичне ПТРЗ, яке визначає умови, за яких коректний алгоритм збереже всі автентичні ребра моделі. Для обґрунтування правил орієнтації ребер потрібні жорсткіші припущення, оскільки необхідно розпізнавати залежність змінних, поєднаних шляхом, тобто послідовністю ребер. Для формулювання прагматичних припущень тестабільності доцільно брати до уваги кількісні характеристики сили залежностей (наприклад, кореляції або величини p -значень за результатами тестування). І ті самі характеристики треба включати в умови правил орієнтації ребер. Але розробники відомих правил абстрагувалися від кількісного аспекту зв'язків і сформулювали правила орієнтації, спираючись тільки на самі факти (не)залежності в локальних конфігураціях. Отже, відповідні припущення також доводиться формувати на концептуальному рівні, суто в термінах фактів (не)залежності.

Умова коректного функціонування правил $\mathcal{R}0$ та $\mathcal{R}1$ формулюється як відповідне припущення тестабільності. Варіанти 2-реберного ланцюга $X \rightarrow Y \rightarrow Z$, $X \leftarrow Y \rightarrow Z$ та $X \leftrightarrow Y \rightarrow Z$ репрезентуємо як $X * \rightarrow Y \rightarrow Z$.

Припущення тестабільності 2-реберної залежності (ПТ2РЗ):

якщо в моделі існує ланцюг $X * \rightarrow Y \rightarrow Z$ і немає $X * \rightarrow Z$, то для всіх наборів $S \in \Lambda_{\text{Emp}}(X, Z)$ таких, що $Y \notin S$, буде значуща залежність $\neg \text{Ind}_{\text{Emp}}(X; S; Z)$, зокрема, буде $\neg \text{Ind}_{\text{Emp}}(X; ; Z)$.

Порушення ПТ2РЗ може призводити, зокрема, до помилкової дії правила $\mathcal{R}0$. Водночас на практиці можливі інші ситуації, коли правило $\mathcal{R}0$ не буде застосоване там, де повинно бути застосоване. Нехай у моделі існують колізори $X * \rightarrow Y \leftarrow * Z$ і деякі ланцюги (ланцюг) між X та Z довжиною чотири чи більше ребер. У цій ситуації можливий невдалий сценарій виведення, який заводить правило $\mathcal{R}0$ орієнтувати колізори $X * \rightarrow Y \leftarrow * Z$. Невдалий сценарій реалізується, якщо алгоритм «наштовхнеться» на «некоректний» сепаратор для X та Z раніше, ніж знайде коректний. Некоректний сепаратор містить змінну Y . Так може статися, коли провокована залежність через колізори $X * \rightarrow Y \leftarrow * Z$ (майже) анігілюється з залежністю через відкриті ланцюги (ланцюг), так що отримаємо $\text{Ind}_{\text{Emp}}(X; S; Z)$, $Y \in S$. Умовою, що невдалий сценарій не відбудеться, є виконання **припущення тестабільності первинної провокованої залежності**:

якщо маємо $X * \rightarrow Y \leftarrow * Z$, то для всіх $S \in \Lambda_{\text{Emp}}(X, Z)$ таких, що $Y \in S$, буде $\neg \text{Ind}_{\text{Emp}}(X; S; Z)$.

Припущення з такою назвою, яке сформульовано в [20], не виконує вказаної функції.

Припущення тестабільності 3-реберної залежності (ПТ3РЗ) формулюється так:

якщо в моделі існує ланцюг $X * \rightarrow Y * \rightarrow Z * \rightarrow W$, то для всіх таких $S \in [\mathbf{K}_{\text{Emp}}(X, W) \cup \{\text{Null}\}]$, які не включають Y або Z ($Y, Z \notin S$), буде $\neg \text{Ind}_{\text{Emp}}(X; S; W)$ (зокрема, виконується $\neg \text{Ind}_{\text{Emp}}(X; ; W)$).

Одне з трьох ребер указанного ланцюга може бути біорієнтованим.

Можна записати ПТЗРЗ в іншій формі:

якщо маємо ланцюг $X \leftrightarrow Y \leftrightarrow Z \leftrightarrow W$, то або $Y \in \text{Sep}_{\min}(X, W)$, або $Z \in \text{Sep}_{\min}(X, W)$, або X та W не сепаруються.

Значення ПТЗРЗ у виведенні моделі можна пояснити на прикладі неанцестральної структури, показаної на рис. 4, а.

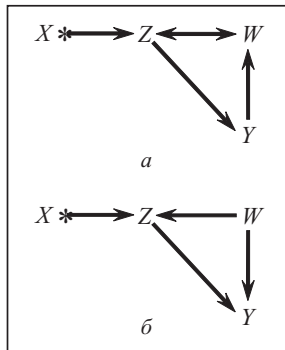


Рис. 4. Приклади структури: неанцестральна (а); анцестральна (б)

Ланцюг $X \leftrightarrow Z \rightarrow Y \rightarrow W$ має забезпечувати безумовну залежність між X та W . Але нехай ПТЗРЗ порушене, так що алгоритм отримує результат тесту $\text{Ind}_{\text{Emp}}(X; W)$. Тоді ілюзорне ребро $X \leftrightarrow W$ вилючається, а правило $\mathcal{R}0$ орієнтує $X \leftrightarrow Z \leftrightarrow W$. Оскільки маємо $\text{Ind}_{\text{Emp}}(X; Z; Y)$, то виконуються умови правила (18), отже, отримаємо $Z \leftrightarrow W \leftrightarrow Y$. Відбулась помилкова орієнтація: замість каузального зв'язку $W \leftarrow Y$ виведено біорієнтований $W \leftrightarrow Y$. (Утім, разом з цією помилкою отримано кілька адекватних орієнтацій.) Натомість, якщо ПТЗРЗ виконується, то результатом виведення в цьому разі буде неорієнтована модель.

Для того щоб запобігти некоректному застосуванню правила (18) для подібних моделей, вказаного припущення недостатньо. У цій ситуації (рис. 4, а) припущення каузальної «не-оманливості» може порушуватися й в іншому форматі тесту. Зокрема, практично можливим є емпіричний факт $\text{Ind}_{\text{Emp}}(X; Y; W)$. Тоді отримаємо той самий помилковий результат. Отже, для обґрунтування коректності правила $\mathcal{R}4_{(1)}^{(N)}$ треба прийняти **припущення тестабільності 1-дистантної провокованої залежності** (ПТ1ДПЗ):

$$X \leftrightarrow Y \leftarrow Z \text{ \& } Y \rightarrow Q \text{ \& } S \in \mathbf{K}_{\text{Emp}}(X, Z) \text{ \& } Q \in S \Rightarrow \neg \text{Ind}_{\text{Emp}}(X; S; Z). \quad (21)$$

В емпіричному сенсі припущення (21) жорсткіше за наведене раніше припущення тестабільності первинної провокованої залежності. Але ще більше підсилювати вимоги недоцільно. Є можливість продуктивної роботи правила $\mathcal{R}0$ у ситуаціях, коли з точки зору d-сепарацій для цього немає умов. Нехай в структурі моделі існують колізор $X \leftrightarrow Y \leftrightarrow Z$ та оршлях $Y \rightarrow Q \rightarrow \dots \rightarrow V \rightarrow Z$. Для пари X, Z не існує графового сепаратора. Але емпіричний сепаратор («нелегітимний», втім, «толерантний») може існувати. Наприклад, може виконуватися $\text{Ind}_{\text{Emp}}(X; V; Z)$. Тоді відкривається можливість орієнтувати колізор за допомогою правила $\mathcal{R}0$.

Розглянемо **припущення тестабільності комбінованої 3-реберної залежності** (ПТКЗРЗ), що формулюється так:

якщо в моделі існує шлях $A \leftrightarrow B \leftarrow D \leftrightarrow C$, причому фрагмент $B \leftarrow D \leftrightarrow C$ є ланцюгом і немає ребра $A \leftrightarrow C$, то для всіх $S \in \mathbf{L}_{\text{Emp}}(A, C)$ таких, що $B \in S, D \notin S$, буде $\neg \text{Ind}_{\text{Emp}}(A; S; C)$, зокрема, виконується $\neg \text{Ind}_{\text{Emp}}(A; B; C)$.

Можна сформулювати звужену версію припущення ПТКЗРЗ: якщо в моделі існує ребро $A \leftrightarrow B$, ланцюг $B \leftarrow D \leftrightarrow C$ та виконується $B \in \text{Sep}_{\min}(A, C)$, то неодмінно D також належить цьому $\text{Sep}_{\min}(A, C)$.

Виконання ПТКЗРЗ необхідне для коректної роботи правила $\mathcal{R}4_{(1)}^{(N)}$. Типова ситуація, де застосовується правило $\mathcal{R}4_{(1)}^{(N)}$, проілюстрована на рис. 4, б. Маємо шлях $X \leftrightarrow Z \leftarrow W \rightarrow Y$. Якщо за умови кондиціонування змінної Z за-

лежність через цей шлях буде незначущою, то отримаємо $\text{Ind}_{\text{Emp}}(X; Z; Y)$. Тоді будуть виконані умови правила (18) і в результаті отримаємо орієнтації $Z \leftrightarrow W \leftrightarrow Y$, тобто помилкові орієнтації двох ребер. Порушення ПТКЗРЗ призвело до того, що спрацював інший випадок правила $\mathcal{R}4_{(1)}^{(N)}$.

Зрозуміло, що запропоновані припущення не охоплюють всіх ризиків емпіричного виведення моделі. Наприклад, слабкою може бути залежність, що утворюється на основі чотирьох ребер. Розглянемо генеративну модель, зображену на рис. 5, а. (У цій структурі існують два майже цикли: $W \rightarrow Z \rightarrow Y \leftrightarrow W$ та $Q \rightarrow W \rightarrow Z \rightarrow Y \leftrightarrow Q$.) Якщо для цієї моделі припущення каузальної «не-оманливості» виконується в повному обсязі, то буде виведена модель, показана на рис. 5, б. (Цей кінцевий результат досягається відразу після застосування правил $\mathcal{R}0$ та $\mathcal{R}1$.) Структура містить ілюзорне ребро $Y \leftarrow V$.

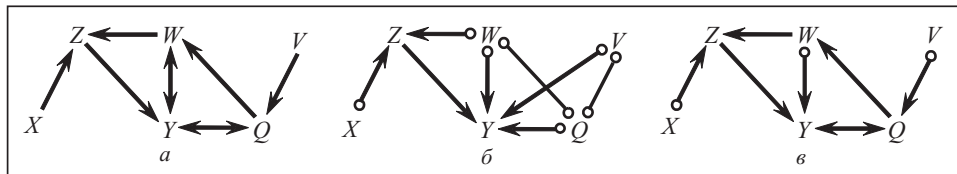


Рис. 5. Приклад виведення неанцестральної моделі: генеративна модель (а); виведена модель (б); проміжна виведена модель (в)

Тепер розглянемо, що станеться, якщо безумовна залежність через ланцюг $Y \leftarrow Z \leftarrow W \leftarrow Q \leftarrow V$ не тестується, тобто маємо $\text{Ind}_{\text{Emp}}(Y;;V)$. (Слабка 4-реберна залежність.) Тоді у виведеній моделі ілюзорного ребра не буде. На етапі 1 алгоритм виведе адекватний набір неорієнтованих ребер, використавши, зокрема, такі результати тестів: $\text{Ind}(V;;Y)$, $\text{Ind}(X;;W)$, $\text{Ind}(V;Q;W)$ та $\text{Ind}(X;\{Z,W\};Y)$. Правило $\mathcal{R}0$ встановить орієнтації ребер $X \circ \rightarrow Z \leftarrow W$, $Y \circ \rightarrow Q \leftarrow V$ та $Z \circ \rightarrow Y \leftarrow Q$. Правило $\mathcal{R}1$ орієнтує ребра $Z \rightarrow Y$, $W \leftarrow Q$ та $Z \leftarrow W$. У цій структурі для ланцюга $W \rightarrow Z \rightarrow Y$ і ребра $W \circ \rightarrow Y$ має бути застосовано правило $\mathcal{R}2$; воно орієнтує ребро $W \circ \rightarrow Y$. Для фрагмента зі змінними V, Q, Y, W правило (18) не спрацює, тому що умови (18) вимагають $Y \circ \rightarrow W$, а маємо $Y \leftarrow W$. Результат виведення відображено на рис. 5, в. Маємо шлях частково невизначеного характеру довжиною у три ребра $V \circ \rightarrow Q \rightarrow W \circ \rightarrow Y$. Якщо уточнити ребро $W \circ \rightarrow Y$ як $W \rightarrow Y$, виникне ланцюг довжиною три ребра, який суперечить факту $\text{Ind}_{\text{Emp}}(Y;;V)$ (з огляду на ПТЗРЗ). Отже, необхідно уточнити ребро як $W \leftrightarrow Y$. Але якими засобами? Для фрагмента зі змінними X, Z, W, Y правило $\mathcal{R}4_{(1)}^{(N)}$ не виконується (до того ж маємо $W \in \text{Sep}_{\min}(X, Y)$, що відповідає випадку (17), який не дає біорієнтованого ребра). Як вже було зазначено, правило (18) не виконується для фрагмента зі змінними V, Q, Y, W через те, що умови вимагають $Y \circ \rightarrow W$. Ця частина умови правила успадкована від оригінального формулювання, призначеного для анцестральних моделей (див. (10) і (12)). Можна послабити цю частину умови, а саме: в формулюванні правила (18) можна $D \circ \rightarrow C$ замінити $D \leftrightarrow C$. Після такої корекції остаточна редакція цього випадку правила $\mathcal{R}4_{(1)}^{(N)}$ має вигляд

$$A \leftrightarrow B \rightarrow C \ \& \ B \leftarrow C \leftrightarrow D \leftrightarrow C \ \& \ D \notin \text{Sep}_{\min}(A, C) \Rightarrow B \leftrightarrow D \leftrightarrow C. \quad (22)$$

Формулювання (22) дасть змогу виконувати орієнтацію ребер у щойно описаній ситуації (що виникає внаслідок слабкості 4-реберних ланцюгів). Але правило (22) призведе до суттєвої проблеми в іншій ситуації, викладеній раніше, коли маємо слабкий 3-реберний ланцюг (коли ПТЗРЗ або ПТКЗРЗ не виконуються). Сценарій такої ситуації, описано раніше та проілюстровано на рис. 3, б. (Докладніше див. [20].) Тоді модель, зображена на рис. 3, б, не буде кінцевим результатом. За тим сценарієм виконуються умови правила (22) і правило ревізує орієнтацію ребра $Z \rightarrow Y$, замінивши її ребром $Z \leftrightarrow Y$. Відбувається «сутичка» правил орієнтації! По суті, застосувати (22) у цій ситуації означає відхилити результат дії правила $\mathfrak{R}1$. Треба шукати інший розв'язок задачі.

ВИСНОВКИ

Досліджено проблеми виведення каузальних моделей із даних у ситуації, коли діють латентні конфаундери (приховані причини, які впливають одночасно на спостережувану причину та її наслідок). Для таких моделей запропоновано корекції правил орієнтації ребер, які забезпечать відтворення адекватних орієнтацій (спрямованості) зв'язків у більшості практично важливих випадків. Як було показано, повнота відтворення зв'язків залежить від: довжини ланцюгів, які забезпечують залежність, що розпізнається тестами; структурного контексту, в який занурено каузальні зв'язки. Серед набору корегованих правил орієнтації ребер тільки одне правило ($\mathfrak{R}1$) виводить каузальний зв'язок.

Сформульовано набір припущень тестабільності залежностей, що описують умови, у яких логічно коректний алгоритм виведе адекватну модель з латентними конфаундерами. Припущення тестабільності залежностей встановлюють вимоги до методів тестування незалежності. Серед запропонованих припущень центральне місце посідає значущість залежностей, утворених на основі 3-реберних шляхів. Сформульовані припущення розраховані на алгоритми, які відшукують локально-мінімальні сепаратори. (Деякі з правил локально-мінімальної сепарації можуть спиратися на додаткові припущення.) Наведені припущення підтримують коректність «прямого» виведення моделі, без перегляду зроблених проміжних висновків. Можна послабити деякі з запропонованих припущень, але тоді виведення моделі не буде «прямим». Тоді для того щоб уникнути помилок, коли алгоритм стикається з якоюсь емпіричною аномалією, доведеться включити в алгоритм додаткові тести, процедури перевірки або навіть етапи виявлення помилок, розплутування суперечностей та ревізії висновків (див., зокрема, [25, 27, 29]).

СПИСОК ЛІТЕРАТУРИ

1. Pearl J. Causality: models, reasoning, and inference. 2nd ed. Cambridge: Cambridge Univ. Press, 2009. 464 p.
2. Spirtes P., Glymour C., Scheines R. Causation, prediction and search. New York: MIT Press, 2001. 543 p.
3. Maathuis M., Drton M., Lauritzen S., Wainwright M. (Eds.). Handbook of Graphical Models. Boca Raton, FL: Chapman & Hall/CRC Press, 2018. 536 p.
4. Glymour C., Zhang K., Spirtes P. Review of causal discovery methods based on graphical models. *Frontiers in Genetics*. 2019. Vol. 10, Article 524. 15 p. <https://doi.org/10.3389/fgene.2019.00524>.

5. Spirtes P., Zhang K. Causal discovery and inference: concepts and recent methodological advances. *Applied Informatics*. 2016. Vol. 3, Iss. 3. <https://doi.org/10.1186/s40535-016-0018-x>.
6. Spirtes P., Meek C., Richardson T. An algorithm for causal inference in the presence of latent variables and selection bias. In: *Computation, Causation, and Discovery*. Glymour C., Cooper G. (Eds.). AAAI Press, Menlo Park, CA. 1999. P. 211–252.
7. Meek C. Causal inference and causal explanation with background knowledge. *Proc. of the 11th Conf. on Uncertainty in Artificial Intelligence*. Besnard P., Hanks S. (Eds.). San Mateo: Morgan Kaufmann Publ, 1995. P. 403–410.
8. Zhang J. On the completeness of orientation rules for causal discovery in the presence of latent confounders and selection bias. *Artificial Intelligence*. 2008. Vol. 172. P. 1873–1896. <https://doi.org/10.1016/j.artint.2008.08.001>.
9. Ali R.A., Richardson T., Spirtes P., Zhang J. Orientation rules for constructing Markov equivalence classes for maximal ancestral graphs. Tech. Rep. 476, Dept. of Statistics, University of Washington, WA. 2005. 41 p.
10. Verma T., Pearl. J. Causal networks: Semantics and expressiveness. *Proc. of 4th Workshop on Uncertainty in Artificial Intelligence*. Minneapolis: Mountain View, 1988. P. 352–359.
11. Meek C. Strong completeness and faithfulness in Bayesian networks. *Proc. of the 11th Conf. on Uncertainty in Artificial Intelligence*. San Mateo: Morgan Kaufmann Publ, 1995. P. 411–418.
12. Balabanov O.S. Induced dependence, factor interaction, and discriminating between causal structures. *Cybernetics and Systems Analysis*. 2016. Vol. 52, N 1. P. 8–19. <https://doi.org/10.1007/s10559-016-9794-5>.
13. Richardson T., Spirtes P. Ancestral graph Markov models. *The Annals of Statistics*. 2002. Vol. 30, N 4. P. 962–1030. <https://doi.org/10.1214/aos/1031689015>.
14. Balabanov A.S. Minimal separators in dependency structures: Properties and identification. *Cybernetics and Systems Analysis*. 2008. Vol. 44, N 6. P. 803–815. <https://doi.org/10.1007/s10559-008-9055-3>.
15. Balabanov A.S. Construction of minimal d-separators in a dependency system. *Cybernetics and Systems Analysis*. 2009. Vol. 45, N 5. P. 703–713. <https://doi.org/10.1007/s10559-009-9136-y>.
16. Balabanov O.S. Logic of minimal separation in causal networks. *Cybernetics and Systems Analysis*. 2013. Vol. 49, N 2. P. 191–200. <https://doi.org/10.1007/s10559-013-9499-y>.
17. Balabanov O.S. Causal nets: Analysis, synthesis and inference from statistical data. Doctor of math. sciences thesis, V.M. Glushkov Institute of Cybernetics, Kyiv, Ukraine, 2014. 353 p. [In Ukrainian].
18. Ali R.A., Richardson T., Spirtes P., Zhang J. Towards characterizing equivalence classes for directed acyclic graphs with latent variables. *Proc. of the 21st Conf. on Uncertainty in Artificial Intelligence*. 2005. P. 10–17.
19. Ali R.A., Richardson T., Spirtes P. Markov equivalence for ancestral graphs. *The Annals of Statistics*. 2009. Vol. 37, N 5B. P. 2808–2837. DOI: 10.1214/08-AOS626.
20. Balabanov O.S. Causal inference from data under presence of latent confounders. Some inadequacy problems (revised). 2021. 18 p. [in Ukrainian]. <https://doi.org/10.13140/RG.2.2.25341.69600>.

21. Tian J., Pearl J. On the testable implications of causal models with hidden variables. *Proc. of the 18th Conf. on Uncertainty in Artificial Intelligence (UAI'02)*. San Francisco: Morgan Kaufmann Publ., 2002. P. 519–527.
22. Shpitser I., Pearl J. Dormant Independence. *Proc. of 23rd AAAI Conf. on Artificial Intelligence (AAAI 2008)*. Chicago, Illinois: AAAI Press, 2008. P. 1081–1087.
23. Shpitser I., Richardson T., Robins J. Testing edges by truncations. *Proc. of the 21st Intern. Joint Conf. on Artificial Intelligence (IJCAI 2009)*. Pasadena, California, 2009. P. 1957–1963.
24. Nowzohour C., Maathuis M.H., Evans R.J., Bühlmann P. Distributional equivalence and structure learning for bow-free acyclic path diagrams. *Electron. J. Statist.* 2017. Vol. 11, N 2. P. 5342–5374. <https://doi.org/10.1214/17-EJS1372>.
25. Ramsey J., Spirtes P., Zhang J. Adjacency-Faithfulness and conservative causal inference. *Proc. of the 22nd Conf. on Uncertainty in Artificial Intelligence*. Oregon: AUAI Press, 2006. P. 401–408.
26. Zhang J., Spirtes P. Detection of unfaithfulness and robust causal inference. *Minds and Machines*. 2008. Vol. 18, N 2. P. 239–271.
27. Spirtes P., Zhang J. A uniformly consistent estimator of causal effects under the k-triangle-faithfulness assumption. *Statist. Sci.* 2014. Vol. 29, N 4. P. 662–678. <https://doi.org/10.1214/13-STS429>.
28. Marx A., Gretton A., Mooij J.M. A weaker faithfulness assumption based on triple interactions. 2021. URL: <https://arxiv.org/abs/2010.14265>.
29. Balabanov A.S. Inference of structures of models of probabilistic dependences from statistical data. *Cybernetics and Systems Analysis*. 2005. Vol. 41, N 6. P. 808–817. <https://doi.org/10.1007/s10559-006-0019-1>.

O.S. Balabanov

LOGIC OF CAUSAL INFERENCE FROM DATA UNDER PRESENCE OF LATENT CONFOUNDERS

Abstract. The problems of causal inference from data (by independence-based methods) when latent confounders are allowed are examined. We demonstrate that the well-known rules of edge orientation may perform wrong under presence of latent confounders. We propose the corrections to the rules aiming to successfully extend them for inference of models beyond the class of ancestral models. The necessary assumptions to justify an adequate model inference from data are suggested.

Keywords: causal network, d-separation, conditional independence, edge orientation rules, confounder, collider, illusive edge, dependence testability assumptions.

Надійшла до редакції 20.09.2021