

XXIII Forum Informatyki Teoretycznej

Zakopane, 24-26 IV 2009

<http://fit.tcs.uj.edu.pl/2009/>

Harmonogram referatów

• Piątek 24 IV

- 13:00 - 15:30 Obiad
- 15:30 - 16:45 Prezentacja referatów
 - * 15:30-15:45 Łukasz Kowalik, Uniwersytet Warszawski, *Wykładowe algorytmy aproksymacyjne dla problemów*
 - * 15:45-16:00 Frantisek Galcik, University of P. J. Safarik, *Communication in radio networks with long-range interference*
 - * 16:00-16:15 Przemysław Gordinowicz, Politechnika Łódzka, *Szybki algorytm wyznaczania maksymalnego przepływu w sieciach.*
 - * 16:15-16:45 Paweł Prałat, Dalhousie University, *Ściganie bandytów w grafie losowym (Chasing robbers on random graphs: zigzag theorem)*
- 16:45 - 17:30 Przerwa kawowa
- 17:30 - 18:45 Prezentacja referatów
 - * 17:30-17:45 Wiktor Żelazny, Uniwersytet Jagielloński, *Szczyśliwe etykietowania grafów*
 - * 17:45-18:00 Sebastian Czerwiński, Uniwersytet Zielonogórski, *Rozgrywane szczyśliwe kolorowanie grafów*
 - * 18:00-18:15 Marcin Piątkowski, Uniwersytet Mikołaja Kopernika, *Zastosowanie grafów podstów w dowodach kombinatorycznych własności słów*
 - * 18:15-18:30 Lech Duraj, Uniwersytet Jagielloński, *Optymalna orientacja grafów on-line*
 - * 18:30-18:45 Artur Jeż, Uniwersytet Wrocławski, *Weryfikacja on-line funkcji π oraz π'*
- 18:45 - 20:30 Kolacja
- 20:30 - 21:45 Prezentacja referatów
 - * 20:30-20:45 Łukasz Jeż, Uniwersytet Wrocławski, *Zrandomizowane szeregowanie pakietów online przeciwko adwersarzowi adaptywnemu*
 - * 20:45-21:00 Jan Jeżabek, Uniwersytet Jagielloński, *Resource Augmentation for QoS Buffer Management with Agreeable Deadlines*
 - * 21:00-21:15 Michał Lasoń, Uniwersytet Jagielloński, *O hipotezie Jaegera*
 - * 21:15-21:30 Bartosz Walczak, Uniwersytet Jagielloński, *O pewnej reprezentacji podstów słowa Fibonacciego*
 - * 21:30-21:45 Tomasz Waleń, Uniwersytet Warszawski, *Powtórzenia w słowach*

• Sobota 25 IV

- 8:00 - 9:00 Śniadanie
- 9:00 - 10:30 Prezentacja referatów
 - * 9:00-9:15 Łukasz Mikulski, Uniwersytet Mikołaja Kopernika, *Aksjomat regularności z punktu widzenia informatyki*

- * 9:15-9:30 Mateusz Kostanek, Uniwersytet Jagielloński, *A limit-colimit coincidence theorem for $\mathcal{Q}$ -dcpo*
- * 9:30-9:45 Piotr Wiśniewski, Uniwersytet Mikołaja Kopernika, *Problem automatycznej dereferencji w językach rodziny SBQL*
- * 9:45-10:00 Szymon Toruńczyk, Uniwersytet Warszawski, *Języki min-regularne*
- * 10:00-10:15 Jakub Michaliszyn, Uniwersytet Wrocławski, *O rozstrzygalności fragmentu strzeżonego z przechodnim domknięciem*
- * 10:15-10:30 Piotr Faliszewski, Akademia Górniczo-Hutnicza, *The Complexity of Power-Index Comparison*
- 10:30 - 14:30 Spacer po górach
- 14:30 - 15:45 Obiad
- 15:45 - 17:15 Prezentacja referatów
 - * 15:45-16:00 Piotr Przymus, Uniwersytet Mikołaja Kopernika, *Rozpoznawanie wzorców sygnałów*
 - * 16:00-16:15 Krzysztof Rykaczewski, Uniwersytet Mikołaja Kopernika, *Metody numeryczne służące do rozpoznawania wzorców pewnych sygnałów*
 - * 16:15-16:30 Michal Mati, University of P. J. Safarik, *Rfam search speed-up filtering algorithm*
 - * 16:30-16:45 Tomasz Idziaszek, Uniwersytet Warszawski, *Języki drzew nieskończonych*
 - * 16:45-17:00 Jerzy Marcinkowski, Jakub Michaliszyn, Uniwersytet Wrocławski, *The cost of being co-Buchi is nonlinear*
 - * 17:00-17:15 Zenon Sadowski, Uniwersytet w Białymstoku, *O problemie istnienia języków zupełnych w klasach semantycznych*
- 17:15 - 17:45 Przerwa kawowa
- 17:45 - 19:00 Prezentacja referatów
 - * 17:45-18:00 Kamila Agata Barylska, Uniwersytet Mikołaja Kopernika, *Problemy dotyczące trwałości w rozszerzeniach p/t-sieci*
 - * 18:00-18:15 Radosław Głowiński, Uniwersytet Mikołaja Kopernika, *Unikanie wzorców w tekstach w dwóch wymiarach*
 - * 18:15-18:30 Sebastian Smoczyński, Uniwersytet Mikołaja Kopernika, *Generowanie drzew binarnych w naturalnym porządku*
 - * 18:30-19:45 Łukasz Pankowski, Uniwersytet Gdański, *Otrzymywanie klucza kryptograficznego z pewnych prawie separowalnych stanów kwantowych o związanym splątaniu*
 - * 18:45-19:00 Jan Otop, Uniwersytet Wrocławski, *O rodzajach E-unifikacji*
- 20:00 Uroczysta kolacja
- Niedziela 26 IV
 - 9:00 - 10:00 Śniadanie
 - 13:00 - 15:00 Obiad
 - Zakończenie FIT

Communication in known topology radio networks with long-range interference

František Galčík*

Abstract

We study communication in known topology radio networks with the presence of interference constraints. A *radio network* is a collection of autonomous stations that are referred to as *nodes*. All network nodes work in globally synchronized time-slots (*rounds*) and operate on the same radio frequency band. In each round, a node works either as a *receiver* or as a *transmitter*.

We consider a real-world situation when transmission of a node produces an interference in the area that is larger than the area where the transmitted message can be received. For each node, there is an area where signal of its transmission has too low intensity to be decoded by a receiver, but is strong enough to interfere with other simultaneously incoming transmissions. Such a setting is modelled by proposed interference reachability graphs that generalize the standard graph model. Let $R_T(v)$, called a *transmission range* of node v , be a set of network nodes that can receive and decode transmission of the node v directly. Further, let $R_I(v)$, called an *interference range* of node v , be a set of network nodes where transmission of the node v can cause interference by its own or together with other simultaneously incoming transmissions. Obviously, $R_T(v) \subseteq R_I(v)$. An interference reachability graph $G = (V, E_T \cup E_I)$ is an undirected (symmetric) graph where:

- V is a set of network nodes,
- E_T is a set of *transmission edges* such that $(u, v) \in E_T \iff v \in R_T(u)$, and
- E_I is a set of *interference edges* such that $(u, v) \in E_I \iff v \in R_I(u) \setminus R_T(u)$.

In a given round, a node u working as a receiver successfully receives a message transmitted by a node v if and only if $u \in R_T(v)$ and there is no other simultaneously transmitting node w such that $u \in R_I(w)$.

One of the most important communication tasks in communication networks is the *broadcasting*. In this task, a message (information) originated in a given source node has to be disseminated in the whole (multi-hop) network. Since we assume that the network topology is known for all nodes, the goal is to design polynomial-time algorithm that computes schedule of transmissions realizing the broadcasting task. The main measure of effectiveness is the length (number of rounds) of the computed schedule.

We give a survey of our recent results related to the broadcasting problem in this model. Focusing on bipartite graphs, we introduce interference selective families as an useful combinatorial tool and show how to construct such families with small size. We compare these results to known results in the standard graph model of radio networks. Finally, we discuss time and energy efficient broadcasting schedules.

Presented results will include results from the paper F. Galčík, *Centralized communication in radio networks with strong interference*, SIROCCO 2008, LNCS 5058, pp. 277-290, and recent joint work with L. Gaşieniec (The University of Liverpool, UK) and A. Lingas (Lund University, Sweden).

*Institute of Computer Science, P.J. Šafárik University in Košice, Jesenná 5, 040 01 Košice, Slovakia. E-mail: frantisek.galcik@upjs.sk

SZYBKI ALGORYTM WYZNACZANIA MAKSYMALNEGO PRZEPŁYWU W SIECIACH.

PRZEMYSŁAW GORDINOWICZ

Wyznaczanie maksymalnego przepływu w sieciach to klasyczny problem optymalizacji dyskretnej z wieloma zastosowaniami. W niniejszym referacie zostanie przedstawiony nowy algorytm działający w oparciu o metodę Dinica[1]. Czas działania algorytmów opartych o tę metodę dla losowych sieci przepływowych jest wyraźnie (o rząd wielkości) niższy niż dla najgorszego przypadku, użytego do analizy złożoności obliczeniowej.

Prezentowane usprawnienie metody Dinica polega na wymianie jednego z jej podprogramów — algorytmu wyznaczającego przepływ nasycający sieć warstwową. W referacie zostanie przedstawiona parametryzowalna i łatwa w implementacji rodzina algorytmów znajdujących przepływ nasycający[4]. Jakkolwiek wspólna dla całej rodziny złożoność obliczeniowa jest słaba $O(V^2L)$, gdzie L jest ilością warstw sieci, jednak dla części wybranych z niej algorytmów łatwo udowodnić złożoność $O(V^2)$. Co więcej, w rodzinie znajduje się szczególnie interesujący „zrównoważony” algorytm, dla którego oszacowana złożoność obliczeniowa to ciągle $O(V^2L)$, lecz nie jest znany przykład sieci warstwowej, która wymagałaby więcej niż $O(E + V^{(3/2)})$ czasu działania. Wszystkie rozważane algorytmy działają szybko (szybciej niż rozwiązania konkurencyjne [2][3]) w eksperymentalnych porównaniach dla losowych przykładów sieci warstwowych.

Wychodząc od przedstawianej rodziny, zostaną skonstruowane trzy algorytmy znajdujące maksymalny przepływ w sieciach w czasie pesymistycznym $O(V^3)$, podczas gdy „zrównoważony” wariant algorytmu jest teoretycznie $O(V^4)$. Jednak czas działania w przypadku eksperymentalnym (losowym) nie przekracza rzędu V^2 .

LITERATURA

- [1] Dinic E. A., Algorithm for solution of a problem of maximal flow in a network with power estimation, *Soviet Math. Dokl.* **11** (1970).
- [2] Malhotra V. M., Pramodh Kumar M., Maheshwari S. N.: *An $O(|V|^3)$ algorithm for finding the maximum flows in network*. Inform. Process. Lett. 7 (1978)
- [3] Sysło M. M., Deo N., Kowalik J. S., *Algorytmy optymalizacji dyskretnej z programami w języku Pascal*, PWN 1995
- [4] P. Gordinowicz, Quick Max-Flow algorithm, *Journal of Mathematical Modeling and Algorithms*, **8** (2009).

E-mail address: pgordin@p.lodz.pl

Chasing robbers on random graphs: zigzag theorem

Paweł Prałat

We study the vertex pursuit game of Cops and Robbers where cops try to capture a robber loose on the vertices of the graph. The minimum number of cops required to win on a given graph G is the cop number of G . We present an asymptotic results for the game of Cops and Robber played on a random graph $G(n,p)$ for a wide range of $p = p(n)$. It has been shown that the cop number as a function of an average degree $d = d(n) = pn$ forms an intriguing zigzag shape.

(joint work with Tomasz Łuczak)

Szczęśliwe etykietowania grafów

Wiktor Żelazny

Przyjmijmy że wierzchołki grafu G poetykietowano liczbami naturalnymi; wagą wierzchołka nazwiemy sumę etykiet na jego sąsiadach. Etykietowanie nazwiemy szczęśliwym jeśli wagi dowolnych dwu sąsiednich wierzchołków są różne. Szczęśliwą liczbą grafu jest najmniejsza liczba k dla której graf ma szczęśliwe etykietowanie liczbami $1..k$. Szczęśliwe etykietowania są naturalnym przeniesieniem na wierzchołki problemu etykietowania krawędzi, pochodzącego od Karońskiego, Łuczaka i Thomasona. Interesuje nas czy szczęśliwa liczba jest ograniczona dla grafów o ograniczonej liczbie chromatycznej; do tej pory nie udało się tego stwierdzić nawet dla grafów dwudzielnych. W moim referacie przedstawię najnowsze rezultaty badań związku między liczbą chromatyczną a szczęśliwą liczbą grafu.

Rozgrywane szczęśliwe kolorowanie grafów

Sebastian Czerwiński

Dla funkcje f określoną na zbiorze wierzchołków grafu G w zbiór $1, 2, \dots, k$, niech $s(u)$ oznacza sumę wartości funkcji f na zbiorze sąsiadów wierzchołka u . Kolorowanie f wierzchołków grafu G nazywamy szczęśliwym, jeżeli dla dowolnej pary sąsiednich wierzchołków u, v mamy, że $s(u)$ jest różne od $s(v)$. W wykładzie pokażemy związek szczęśliwego kolorowania grafów dwudzielnych z pewną grą w zapalanie lampek na grafie.

Zastosowanie grafów podsłów w dowodach kombinatorycznych własności słów

Marcin Piątkowski*

Grafem podsłów słowa w nazywamy minimalny deterministyczny automat skończony akceptujący zbiór wszystkich sufiksów tego słowa. Pozwala on w zwarty sposób reprezentować zbiór wszystkich podsłów w . Najważniejszą zaletą grafu podsłów jest jego liniowy rozmiar względem długości słowa, podczas gdy ilość różnych podsłów może być rzędu $O(|w|^2)$.

Standardowe słowa Sturma to klasa silnie kompresowalnych słów nad alfabetem $\Sigma = \{a, b\}$. Definiowane są one rekurencyjnie. Dla ciągu liczb całkowitych $d = (d_0, d_1, \dots, d_{n-1})$, spełniającego warunki: $d_0 \geq 0$ oraz $d_i > 0$ dla $0 < i < n$, określamy n -te słowo standardowe Sturma za pomocą rekurencji:

$$\begin{aligned}x_{-1} &= b \\x_0 &= a \\x_i &= x_{i-1}^{d_{i-1}} x_{i-2} \quad \text{dla } 0 < i < n\end{aligned}$$

Grafy podsłów standardowych słów Sturma mają bardzo szczególną strukturę implikującą wiele ciekawych własności kombinatorycznych. Jej dokładna analiza pozwala również uprościć dowody wielu znanych faktów opartych wyłącznie na strukturze słowa.

Literatura

- [1] P. Baturó, M. Piątkowski, W. Rytter *Usefulness of Directed Acyclic Subword Graph in Problems Related to Standard Sturmian Words*, Proceedings of Prague Stringology Conference 2008, pp. 193-207.
- [2] M. Lothaire *Applied combinatorics on words*, Cambridge University Press, 2005.

*Badania wspierane przez Ministerstwo Nauki i Szkolnictwa Wyższego, grant N N206 258035

OPTYMALNA ORIENTACJA GRAFÓW ON-LINE

Lech Duraj

Problemy optymalnej orientacji grafu formułuje się następująco: mając dany graf nieskierowany $G = (V, E)$ należy każdej jego krawędzi nadać kierunek tak, aby maksymalizować jedną z dwóch wielkości:

- liczbę par wierzchołków $(u, v) \in V^2$, dla których istnieje ścieżka skierowana z u do v (CONNECTED PAIRS OPTIMAL ORIENTATION)
- sumę stopni spójności wierzchołkowej lub krawędziowej po wszystkich parach wierzchołków (AVERAGE CONNECTIVITY OPTIMAL ORIENTATION)

Oba te problemy można rozważać w wersji off-line (klasycznej), jak również w wersji on-line: dane wejściowe nie muszą być znane algorytmowi od początku w całości, lecz są podawane w turach – najczęściej zakłada się, że w każdej turze dopisywany jest jeden wierzchołek i wychodzące z niego krawędzie do wierzchołków poprzednich. Algorytm musi wszystkie podane krawędzie od razu zorientować, a raz podjęte decyzje są permanentne.

Wersję on-line opisuje się też jako grę pomiędzy dwoma graczami: podającym wierzchołki *Spoilerem*, oraz orientującym krawędzie *Algorytmem*. Celem Spoilera jest ułożenie takiego grafu, aby Algorytm osiągnął na nim wynik możliwie daleki od optymalnego.

Znane jest rozwiązanie problemu CONNECTED PAIRS w wersji off-line – kwadratowy algorytm dający optymalną orientację, która zawsze gwarantuje $\Theta(n^2)$ połączonych par. W moim referacie chciałbym się skupić na rozwiązaniu problemów optymalnej orientacji w wersji on-line. Jako że praktycznie nieopłacalne dla Spoilera jest podawanie grafów nie będących drzewami, rozwiązanie stosuje się do obu problemów i na ogół wystarczy rozważać tylko CONNECTED PAIRS.

Okazuje się, że pewien zachłanny algorytm zapewnia sobie wynik $\Omega(n \frac{\log n}{\log \log n})$ połączonych par dla każdego grafu o n wierzchołkach podawanego on-line. Z drugiej strony, istnieje pewien graf G i strategia jego podawania, która ogranicza wynik dowolnego algorytmu do $O(n \frac{\log n}{\log \log n})$. Łącznie oznacza to, że przy optymalnej grze obu graczy osiągnięty zostanie wynik $\Theta(n \frac{\log n}{\log \log n})$, znacząco mniejszy od off-line'owego $\Theta(n^2)$. Rozważę też wersję problemu z dodatkowymi założeniami o podawanym grafie (takimi jak ograniczony stopień wierzchołka lub minimalny stopień spójności grafu), które dają Algorytmowi dodatkowe informacje i tym samym mogą zwiększyć wynik gry.

Weryfikacja on-line funkcji π oraz π'

Paweł Gawrychowski, Artur Jeż i Łukasz Jeż

Instytut Informatyki Uniwersytet Wrocławski

Niech π_w oznacza funkcję prefiksową używaną w algorytmie Morrisa-Pratta dla słowa w . Rozpatrujemy następujący problem: dla zadanej na wejściu tablicy $A[1..n]$, czy istnieje słowo w , takie że $A = \pi_w$? Jeśli tak, jaki jest minimalny rozmiar alfabetu, dla którego takie słowo istnieje? Prezentujemy algorytm on-line o opóźnieniu stałym, który zwraca słowo realizujący powyższe cele.

Następnie zajmujemy się analogicznym problemem dla funkcji π' , prezentujemy algorytm on-line działający w czasie $\mathcal{O}(n \log n)$.

Randomized Algorithms for Buffer Management with 2-Bounded Delay

Marcin Bienkowski¹, Marek Chrobak², and Łukasz Jeż¹

¹ Instytut Informatyki Uniwersytet Wrocławski *

² Department of Computer Science, University of California, Riverside **

Streszczenie: Problem *Buffer Management with Bounded Delay* to w istocie problem szeregowania online zadań o jednostkowych długościach i dowolnych wagach oraz terminach („deadline’ach”) na jednej maszynie: pakiety o wagach (priorytetach) pojawiają się w dowolnym momencie w przełączniku sieciowym, z którego w każdym kroku przesłać dalej można tylko jeden pakiet. Zadaniem jest zmaksymalizowanie sumy wag zadań, które udało się przesłać, nim upłynął ich termin.

Najlepsze znane ograniczenia dolne na współczynnik konkurencyjności dla tego problemu uzyskiwano przez analizę tzw. instancji 2-ograniczonych, tj. takich w których dla dowolnego pakietu j różnica między jego terminem d_j a czasem zwolnienia r_j nie przekracza 2. Innymi słowy, takich, w których każdy pakiet musi zostać przekazany tuż po zwolnieniu, lub ewentualnie w następnym kroku, a w przeciwnym razie jest tracony.

Dla tego typu instancji optymalny współczynnik konkurencyjności wśród algorytmów deterministycznych wynosi $\phi \approx 1.618$. W przypadku algorytmów zrandomizowanych, optymalny współczynnik przeciwko nieświadomemu adwersarzowi (*oblivious*) wynosi 1.25. Nieznany pozostawał optymalny współczynnik przeciwko adwersarzowi adaptywnemu (*adaptive online*). Dowodzimy, że w tym wypadku optymalnym współczynnikiem jest $4/3$. Dodatkowo, przez drobną modyfikację konstrukcji z ograniczenia dolnego, uzyskujemy ograniczenie 1.2 dla tzw. instancji 2-jednorodnych, tj. takich, w których dla każdego zadania różnica między terminem a czasem zwolnienia wynosi *dokładnie* 2. Tego typu instancje również były badane wcześniej, jako najbardziej ograniczony nietrywialny wariant problemu.

* Supported by MNiSW grants number N206 001 31/0436, 2006–2008 and N N206 1723 33, 2007–2010.

** Supported by NSF grants OISE-0340752 and CCF-0729071.

Resource Augmentation for QoS Buffer Management with Agreeable Deadlines

Jan Jezabek

Theoretical Computer Science Department
Jagiellonian University, Krakow, Poland
`jezabek@tcs.uj.edu.pl`

April 11, 2009

Abstract

We study the following problem arising in network switches: At each step new packets are received, each characterized by its weight (or value) and its deadline. At each step a fixed number of packets whose deadline has not yet passed may be transmitted, while the untransmitted ones are stored in an unbounded buffer. The goal is to maximize the sum of the weights of all the packets transmitted by the switch. In the agreeable deadlines setting if one packet arrives before another one then the first packet's deadline may not be later than that of the second packet.

In this paper we study the effects of resource augmentation on this problem in the agreeable deadlines setting. Specifically we present an on-line algorithm for a switch transmitting two packets per step which is guaranteed to perform at least as well as the optimal off-line algorithm for a switch transmitting one packet per step. This complements our earlier result which showed that in the general setting such an algorithm does not exist.

Key words: buffer management, on-line algorithms, scheduling

O pewnej reprezentacji podsłów słowa Fibonacciego

Bartosz Walczak

Zdefiniuję pewną prostą reprezentację podsłów nieskończonego słowa Fibonacciego za pomocą konkatencji skończonych słów Fibonacciego i udowodnię, że reprezentacja ta charakteryzuje te podsłowa w sposób jednoznaczny. Jako natychmiastowy wniosek otrzymamy liniowy algorytm testowania, czy dane słowo jest podsłowem słowa Fibonacciego, oraz pełny opis pozycji wystąpień danego słowa jako podsłowa w słowie Fibonacciego.

On the Maximal Number of Cubic Subwords in a String

**Marcin Kubica¹, Jakub Radoszewski¹, Wojciech Rytter^{1,2}, and
Tomasz Waleń¹**

¹ Department of Mathematics, Computer Science and Mechanics,
University of Warsaw, Warsaw, Poland
[kubica,jrad,rytter,walen]@mimuw.edu.pl

² Faculty of Mathematics and Informatics,
Copernicus University, Toruń, Poland

We investigate the problem of the maximum number of cubic subwords (of the form www) in a given word. We also consider square subwords (of the form ww). The problem of the maximum number of squares in a word is not well understood. Several new results related to this problem are produced in the paper. We consider two simple problems related to the maximum number of subwords which are squares or which are highly repetitive; then we provide a nontrivial estimation for the number of cubes. We show that the maximum number of squares xx such that x is not a primitive word (nonprimitive squares) in a word of length n is exactly $\lfloor \frac{n}{2} \rfloor - 1$, and the maximum number of subwords of the form x^k , for $k \geq 3$, is exactly $n - 2$. In particular, the maximum number of cubes in a word is not greater than $n - 2$ either. Using very technical properties of occurrences of cubes, we improve this bound significantly. We show that the maximum number of cubes in a word of length n is between $\frac{45}{100}n$ and $\frac{4}{5}n$.

Aksjomat regularności z punktu widzenia informatyki¹

Łukasz Mikulski, Uniwersytet Mikołaja Kopernika w Toruniu

Powstałe na początku XX wieku antynomie (paradoks Russela), związane ze zbyt intuicyjnym i niesformalizowanym podejściem do pojęcia zbioru, zaowocowały rozwojem aksjomatycznej teorii mnogości. Wiele uznanych podręczników akademickich [4,6] ogranicza się do podania aksjomatów algebry zbiorów rozszerzonych o aksjomaty podzbiorów, zbioru potęgowego oraz wyboru nazywając je wyrażającymi „te wszystkie własności pojęcia zbioru, [...] które – na ogół biorąc – obejmują zastosowania teorii mnogości do innych działów matematyki”. Jednakże ostatni z przytoczonych powyżej aksjomatów budzi pewne kontrowersje w środowisku matematycznym.

W dalszych rozdziałach „Wstępu do teorii mnogości” Kuratowskiego, do tej listy podstawowych aksjomatów dołączają aksjomaty nieskończoności oraz zastępowania. Oprócz tych aksjomatów, w pozycjach o nieco ogólniejszym charakterze [3,5], można znaleźć wzmianki o aksjomacie regularności, zwanym też aksjomatem ufundowania. O ile aksjomat ten został przyjęty przez Chronowskiego, to Kuratowski i Mostowski piszą tylko, że „w niektórych wykładach teorii mnogości przyjmowany bywa tzw. aksjomat regularności”, nazywając jednak opisywaną przez niego własność bardzo intuicyjną.

Aksjomat regularności, uniemożliwiający należenie zbioru do samego siebie, ma postać:

$$X \neq \emptyset \Rightarrow \exists_{x \in X} x \cap X = \emptyset \text{ [5]},$$

czyli każdy zbiór niepusty posiada element z nim rozłączny. Jego niezależność od pozostałych aksjomatów udowodnił w roku 1941 roku Bernays (dowód został opublikowany w roku 1954 [2]). Warto też dodać, że paradoks Russela, który był prawdopodobnie jednym z powodów dołączenia do listy aksjomatów teorii mnogości aksjomatu regularności, dowodzi w połączeniu z aksjomatem podzbiorów, że klasa wszystkich zbiorów nie jest zbiorem.

Odrzucenie aksjomatu regularności, lub przyjęcie jego zaprzeczenia, znacznie upraszcza wiele formalnych notacji w informatyce teoretycznej. Każdy graf skierowany z wyróżnionym źródłem odpowiada wówczas pewnemu zbiorowi w modelu teorii mnogości. Zbędne stają się funkcje przejścia w teorii automatów czy systemów tranzycyjnych. Zawarte w tych funkcjach idee można przenieść wprost do definicji stanów, co jest działaniem intuicyjnym i naturalnym. Łatwiej też wyrazić idee rekurencji czy wskaźników. Cytując Aczela, upraszczając takie konstrukcje i dając im elegancki model teoriomnogościowy musimy się liczyć z „karą”, jaką jest pojawienie się zbiorów źle ufundowanych. Jeśli jednak przyjmiemy aksjomat antyregularności, to nie ma już żadnej kary [1].

[1] P. Aczel, Non-Well-Founded Sets, CSLI Lecture Notes, vol. 14, 1988.

[2] P. Bernays, A System of Axiomatic Set Theory, Journal of Symbolic Logic, vol. 19, pp. 81-96, 1954.

[3] A. Chronowski, Elementy teorii mnogości, Wydawnictwo Naukowe Akademii Pedagogicznej, Kraków 2000.

[4] K. Kuratowski, Wstęp do teorii mnogości i topologii, Biblioteka matematyczna, tom 9, PWN, Warszawa 1972.

[5] K. Kuratowski, A. Mostowski, Teoria mnogości, Monografie matematyczne, tom 27, PWN, Warszawa 1966.

[6] H. Rasiowa, Wstęp do matematyki współczesnej, Biblioteka matematyczna, tom 30, PWN, Warszawa 1977.

¹ Badania wspierane przez Ministerstwo Nauki i Szkolnictwa Wyższego, grant N N206 258035

A limit-colimit coincidence theorem for $\mathcal{Q}$ -dcpo

Mateusz Kostanek

One of the widely accepted research methods in mathematics is lifting the essence of your favourite theory to a more general setting. In this talk we want to establish a fact that a significant part of domain theory (our favourite) can be generalized to the theory of $\mathcal{Q}$ -categories. (Equivalently, we can say that we are going to show that a significant part of the theory of $\mathcal{Q}$ -categories comes from generalizing domain-theoretic constructions.) In this way we can obtain some yet unknown results about $\mathcal{Q}$ -categories. For example we prove: (a) a limit-colimit coincidence theorem for $\mathcal{Q}$ -posets; (b) the necessary and sufficient condition for cartesian closedness of the category of complete $\mathcal{Q}$ -posets and Scott-continuous maps.

Problem automatycznej dereferencji w językach rodziny SBQL

Piotr Wiśniewski, Marta Burzańska

W roku 1993 została przedstawiona koncepcja podejścia stosowego w bazach danych. Podejście to wprowadza mechanizmy dobrze znane z języków programowania – stos środowisk i stos rezultatów – do konstrukcji języka zapytań. Dzięki temu konstruowany język wprost można rozszerzyć o konstrukcje niezbędne w językach programowania. W wyniku tych prac opracowano język SBQL. Autorzy niniejszej pracy obserwując rozwój wiedzy o językach zapytań do obiektowych baz danych opartych o podejście stosowe oraz obserwując wzrost popularności języka Python ze względu na jego łatwość i porządek powstałego kodu, rozpoczęli prace koncepcyjne nad językiem PySBQL, łączącym zalety Pythona, jako języka programowania z naturalnością SQBLa jako języka zapytań. Język SBQL przyjmuje założenie, że podstawowym wynikiem zwracanym w wyniku ewaluacji identyfikatora jest referencja – lub kolekcja referencji. Naturalnym efektem ubocznym tej koncepcji jest sytuacja, w której programista naturalnie oczekuje wartości a nie referencji. W odpowiedzi na te sytuacje wprowadzono operację naturalnej dereferencji i konwencję, kiedy dereferencja ta jest wykonana. Dla sytuacji wyjątkowych – odbiegających od konwencji konstruktorzy SQBLa wprowadzili operatory odpowiednio deref i ref. W opinii autorów konwencja ta prowadzi do wielu nieporozumień. Celem referatu jest przedstawienie alternatywnego podejścia do ewaluacji identyfikatora opracowanego przez autorów dla języka PySBQL, a opartego o koncepcje lewej i prawej wartości, dobrze znane z konstrukcji języków programowania.

Języki *min*-regularne

Szymon Toruńczyk*
Uniwersytet Warszawski

Rozpatrujemy nową klasę języków słów nieskończonych, zwanych *min-regularnymi*, które poszerzają klasę języków ω -regularnych. Klasa ta ma dwa równoważne opisy: przez pewien model automatów deterministycznych wyposażonych w liczniki, oraz w języku logiki (słabej logiki monadycznej poszerzonej o pewien nowy kwantyfikator).

Problem pustości języka z tej klasy jest rostrzygalny. Sprowadza się on do tzw. problemu *skończonego przekroju* dla półgrupy macierzy nad *pierścieniem tropikalnym* (jest to zbiór liczb naturalnych wyposażony w działania $(\min, +)$).

*wspólna praca z Mikołajem Bojańczykiem

Decidability of the Guarded Fragment with the Transitive Closure*

Jakub Michaliszyn

April 14, 2009

The guarded fragment of first-order logic, GF, introduced in [1], is a well-known generalisation of modal logics. The main idea is to allow only a restricted form of quantification, simulating local nature of the modal operators $\Diamond$, $\Box$. GF retains a lot of nice properties of modal logics, including (a generalisation of) the tree model property, the finite modal property and the robust decidability. This makes it a promising starting point for logics for reasoning about programs and hardware.

Many extensions and variants of GF have been extensively investigated last years. One important direction is considering satisfiability of GF over restricted classes of structures, in which some distinguished binary symbols are interpreted as transitive relations (or, alternatively phrased, satisfiability of GF extended by positive statements about transitivity of some binary relations). It appeared ([5], [2]) that allowing transitive symbols to appear in arbitrary positions in formulas leads quickly to undecidability. However, if we restrict the usage of transitive symbols to guards only, the satisfiability problem becomes decidable and 2EXPTIME-complete ([7], [6]). The lower bound can be proved even in the presence of only two variables. This two-variable version, $[GF^2 + TG]$, captures (and non-trivially extends) modal logics K4, S4 and S5.

Transitivity of some binary relations is a desirable property in many reasoning tasks. However, to reason about programs it would be nice to have some way of expressing recursion. One idea is to extend GF by least and greatest fixed point operators. It was done in [3]. This way we obtain a powerful logic embedding modal μ -calculus with backward modalities. Another idea is to add to GF some form of the transitive closure operator. In this paper we consider $[GF^{2++}]$, an extension of the two-variable guarded fragment, in which the transitive closure operator can be applied to some atomic formulas. We note that augmenting fragments of first-order logic even by a weak form of the transitive closure leads quickly to undecidability (see, e. g., [4]). Thus we have to be careful: analogously to the variant with transitive relations we restrict the usage of the transitive closure operator $^+$ to symbols appearing only in guards. More formally, the signature has a distinguished subset Σ_+ , containing binary symbols, which may be used only in guards, either individually or under the transitive closure operator. We note that a variant, in which symbols from Σ_+ are used additionally outside guards, but not under the transitive closure, is undecidable.

$[GF^{2++}]$ easily simulates $[GF^2 + TG]$: in a formula of $[GF^2 + TG]$ instead of a transitive relation T we can simply use a transitive closure of a relation T' . However $[GF^{2++}]$ is strictly more expressive than $[GF^2 + TG]$. For example, consider the formula $\exists x(S(x) \wedge \exists y(xT^+y \wedge R(y)))$, stating that from some point in S there exists a T -path to a point in R . This is clearly not a first-order property and thus it cannot be expressed in $[GF^2 + TG]$.

We prove that the satisfiability problem for $[GF^{2++}]$ is decidable in 2EXPTIME, exactly as $[GF^2 + TG]$. Similarly to the case of $[GF^2 + TG]$ the proof is based on a tree-like model property. However, there are some serious complications, due to the fact that $[GF^{2++}]$ can speak both about direct successors and about reachable elements. For example, in contrast to $[GF^2 + TG]$, for a given element its witnesses cannot always be its direct successors.

References

- [1] H. Andreka, I. Nemeti, J. van Benthem. Modal languages and bounded fragments of predicate logic. *Journal of Philosophical Logic*, 27:217-274, 1998.
- [2] H. Ganzinger, C. Meyer, and M. Veanes. The two-variable guarded fragment with transitive relations. In *Proc. LICS'99. IEEE Computer Soc. Press*, 1999.
- [3] E. Grädel, I. Walukiewicz. Guarded Fixed Point Logic. In *Proceedings of 14th IEEE Symposium on Logic in Computer Science LICS '99*, Trento, 1999.
- [4] E. Grädel, M. Otto, E. Rosen. Undecidability results on two-variable logics. *Lecture Notes in Computer Science Nr. 1200*, Springer, pp. 249-260, 1997.
- [5] E. Grädel. On the restraining power of guards. *Journal of Symbolic Logic*, 64(4):1719-1742, 1999.
- [6] E. Kieroński, The two-variable guarded fragment with transitive guards is 2EXPTIME - Hard. *Lecture Notes in Computer Science*, vol. 2620, pp. 299-312, 2003.
- [7] W. Szwaś, L. Tendera. The guarded fragment with transitive guards. *Annals of Pure and Applied Logic*, 128, 227-276, 2004.

*Part of the M.Sc. thesis, under the supervision of Emanuel Kieroński

The Complexity of Power-Index Comparison*

Piotr Faliszewski

Department of Computer Science
AGH University of Science and Technology
Kraków, Poland

Lane A. Hemaspaandra

Department of Computer Science
University of Rochester
Rochester, NY, USA

An n -player weighted voting game is a sequence of n nonnegative integer weights, $w_1, \dots, w_n$, together with a quota q . We denote it as $(w_1, \dots, w_n; q)$. We refer to the player with weight w_i as the i 'th player. Weighted voting games model the following scenario: The players are given a yes/no question (e.g., should we lower the taxes? should we buy out our competitors?) and each player either agrees (answers *yes*) or disagrees (answers *no*). If the total weight of the voters who agree is at least as high as the quota then the result of the game is *yes* (and we say the coalition is *successful*) and otherwise it is *no* (and the coalition is not successful). We define $\text{succ}_G(S)$ to be 1 if S is a successful coalition for G and to be 0 otherwise.

When analyzing weighted voting games it is standard to measure players' power using, e.g., the Shapley-Shubik power index [SS54] or the Banzhaf power index [Ban65]. In essence, these power indices measure the probability that, assuming some coalition formation model, our designated player is critical for the forming coalition. By *critical* we mean here that the coalition is successful with our designated player but is not successful without him or her.

Let $G = (w_1, \dots, w_n; q)$ be a voting game, let i be a player in this game, and let $N = \{1, \dots, n\}$ be the set of all players of G . The Shapley-Shubik power index of player i in game G is defined as $\text{SS}(G, i) = \frac{\text{SS}^*(G, i)}{n!}$, where $\text{SS}^*(G, i)$ is the raw version of the index, $\text{SS}^*(G, i) = \sum_{S \subseteq N - \{i\}} \|S\|!(n - \|S\| - 1)!(\text{succ}_G(S \cup \{i\}) - \text{succ}_G(S))$. The Banzhaf power-index is defined similarly, but it uses a different coalition-formation model.

We study the complexity of the following problem: Given two weighted voting games G' and G'' that each contain a player p , in which of these games is p 's power index value higher? We study this problem with respect to both the Shapley-Shubik power index [SS54] and the Banzhaf power index [Ban65]. Our main result is that for both of these power indices the problem is complete for probabilistic polynomial time (i.e., is PP-complete). We also study the complexity of the raw Shapley-Shubik power index. Deng and Papadimitriou [DP94] showed that the raw Shapley-Shubik power index is #P-metric-complete. We strengthen this by showing that the raw Shapley-Shubik power index is many-one complete for #P. And our strengthening cannot possibly be further improved to parsimonious completeness, since we observe that, in contrast with the raw Banzhaf power index, the raw Shapley-Shubik power index is not #P-parsimonious-complete.

References

- [Ban65] J. Banzhaf. Weighted voting doesn't work: A mathematical analysis. *Rutgers Law Review*, 19:317–343, 1965.
- [DP94] X. Deng and C. Papadimitriou. On the complexity of comparative solution concepts. *Mathematics of Operations Research*, 19(2):257–266, 1994.
- [SS54] L. Shapley and M. Shubik. A method of evaluating the distribution of power in a committee system. *American Political Science Review*, 48:787–792, 1954.

*The first author was partially supported by grant NSF-CCF-0426761. The second author was partially supported by grant NSF-CCF-0426761, a TransCoop grant, and a Friedrich Wilhelm Bessel Research Award.

Piotr Przymus
FIT 2009

“Rozpoznawanie wzorców sygnałów”

Przedmiotem referatu jest zaprezentowanie praktycznej metody wykrywania charakterystycznych wzorców w określonym sygnale. Analizowany sygnał jest zapisem zjawiska biologicznego i otrzymany został w wyniku przeprowadzenia pewnego eksperymentu w zakładzie Hydrobiologii UMK.

W wyniku owego eksperymentu otrzymano dyskretne sygnały, które następnie zanalizowano pod kątem możliwości diagnozowania zjawiska, jego powtarzalności oraz adekwatności do określonych rezultatów badawczych. Pozwoliło to na wyodrębnienie właściwej informacji zawartej w badanych sygnałach, która pozwala na odnalezienie analogii pomiędzy zachodzącym zjawiskiem a cechami zmiennymi środowiska, w którym dokonywane są pomiary. W sygnale można odnaleźć charakterystyczne wzorce, które są współzmiennicze z momentami zachodzenia badanych zjawisk.

W swojej prezentacji przedstawię algorytm wykrywania wzorców sygnału, będących odzwierciedleniem pewnego stanu fizycznego. Powstał on w wyniku mojej współpracy z Krzysztofem Rykaczewskim. Rozpoznane przez algorytm wzorce zostaną poddane dalszej analizie w celu identyfikacji zachodzącego procesu biologicznego przy użyciu sieci neuronowych.

Tworzenie owego algorytmu poprzedzone zostało przeanalizowaniem szeregu metod numerycznych pod kątem przydatności w tym konkretnym zagadnieniu. Szczegółowe omówienie zastosowanych i nie wykorzystanych metod numerycznych, ich skuteczności a także kwestie dotyczące wpływu parametrów na skuteczność prezentowanych metod zostaną zaprezentowane w referacie Krzysztofa Rykaczewskiego.

Niektóre z metod użytych w algorytmie:

- Filtry liniowe i nieliniowe do przetworzenia próbek sygnału.
- Analiza statystyczna sygnału.
- Analiza falkowa do wykrywania charakterystycznych zaburzeń sygnału.
- Transformata Hilberta służąca do wykrycia obwiedni sygnału.

W dalszej części referatu zaprezentuję przykładowe fragmenty badanych sygnałów i rezultaty otrzymane po zastosowaniu opisanego algorytmu. Postaram się w ten sposób wykazać jego skuteczność w poszczególnych aspektach. Omówiona zostanie również implementacja algorytmu i użyte biblioteki.

Krzysztof Rykaczewski

FIT2009

„Metody numeryczne służące do rozpoznawania wzorców pewnych sygnałów”

Referat i prezentacja dotyczyć będą zastosowania modelowania matematycznego i symulacji komputerowej do analizy sygnału pochodzącego od pewnego biologicznego zjawiska dynamicznego.

Uzyskane metodami empirycznymi w pewnym eksperymencie biologicznym dyskretne sygnały poddawane zostały analizie w celu określenia ich diagnostycznych własności i relacji pomiędzy sygnałem a stanem rzeczywistym. Stan rzeczywisty można opisać za pomocą wielu czynników, które grupuje się na procesy ze względu na wyróżnione cechy. Cechy te w jednoznaczny sposób odpowiadają charakterowi wykresu tego procesu. Wstępne eksperymenty oraz dotychczasowe wyniki prac z P. Przymusem wykazują, że takie sformułowanie problemu prowadzi do wyodrębnienia właściwej informacji zawartej w rozpatrywanym sygnale, co pozwala sądzić, że przedstawione wyniki symulacji dają pełną gwarancję efektywnego rozwiązania postawionego problemu. W sygnale daje się znaleźć charakterystyczne kształty, które są współzmiennicze z momentami zachodzenia wspomnianych wyżej procesów. Zmiany jakościowe mogą być w tym wypadku stosowane do detekcji cech środowiska, w którym dokonuje się pomiar.

Do oceny stanu fizycznego na podstawie wykresu można wykorzystać wiele różnych metod analizy sygnałów. W referacie omówiony zostanie przegląd i porównanie najpopularniejszych metod stosowanych w diagnostyce zorientowanej na wyróżnienie pewnych faz w badanych próbkach. Techniki te mogą być stosowane na etapie monitorowania systemu, badań identyfikacyjnych oraz jako standardowa procedura w zaawansowanych systemach kontrolnych. Istotną częścią prezentowanej koncepcji jest automatyzacja procesu weryfikacji zachowań, dzięki czemu będzie możliwa identyfikacja przy pomocy wielowarstwowych sieci neuronowych. Podkreślono problemy z interpretacją wyników analizy sygnału z danych badanych przez autora. Szczególna uwaga została poświęcona zastosowanym matematycznym metodom numerycznym. Wyeksponowano zalety analizy falkowej w procesie identyfikacji struktury częstotliwościowej sygnału, lokalizacji charakterystycznych zaburzeń, dyskryminacji przypadków prawidłowych oraz w wykrywaniu faz nieprawidłowości. Wspomnę także o stosowanych przez nasz zespół filtrach liniowych i nieliniowych oraz metodzie transformaty Hilberta, służącej do wykrycia obwiedni sygnału.

Omówione zostanie zagadnienie wpływu różnych parametrów na jakość analizy oraz testów oceny poprawności wychwytywania ramek. Przedstawione zostaną właściwości analizy wielorozdzielczej z wykorzystaniem różnych falek oraz opisano zastosowany algorytm. Przedyskutowany zostanie problem doboru parametrów analizy do ciągu próbek wartości chwilowych w aspekcie każdej z metod.

Zagadnienie projektowania takich systemów analizy z wykorzystaniem bibliotek analiz sygnałowych oraz problem implementacji algorytmów przetwarzania sygnałów przetwarzanych, będący częścią tego projektu, w aspekcie oprogramowania użytkowego przedyskutowany zostanie w referacie P. Przymusa.

Na koniec przedstawię wnioski, jakie nasunęły się w trakcie prac nad projektem. Dodatkowo proponuję pewne możliwości usprawnień, które mamy nadzieję zastosować w późniejszym czasie.

Rfam search speed-up filtering algorithm

Michal Mati

For the general problem of defining RNA families and finding new members, two methods exist that require modest manual work per family: Covariance models and ERPIN. ERPIN sometimes requires the user to specify score thresholds for each helix, thus requiring more expert input and/or compromising accuracy. This paper focuses on the impractical speed of CMs without sacrificing their accuracy. CMs are used in the Rfam Database to annotate an 8-gigabase genome database called RFAMSEQ for over 100 ncRNA (non-coding RNA) families. CMs are too slow to be used directly: e.g., searching RFAMSEQ to find tRNAs would take about 1 year on a 2.8 GHz Intel Pentium 4 PC. Obviously, this is improving as computers get faster, but there are over 100 families. Moreover, new families continue to be found and sequence data expands rapidly. We propose the filtering algorithm for speeding-up Rfam database search which parallels the original algorithm of Sun and Buhler for profile HMM search.

Języki drzew nieskończonych

Tomasz Idziaszek

Rozstrzygalną charakteryzacją dla logiki L nazywamy algorytm, który dla danego na wejściu automatu stwierdza, czy język rozpoznawany przez ten automat może być zdefiniowany formułą w logice L . Takie charakteryzacje ukazują nam ciekawe związki pomiędzy teorią automatów i logiką, a do ich znajdowania często stosuje się metody algebraiczne.

Pokażę jak można zdefiniować algebrę dla drzew nieskończonych, która posłuży nam w sformułowaniu rozstrzygalnej charakteryzacji dla prostej logiki temporalnej EF.

The cost of being co-Büchi is nonlinear

Jerzy Marcinkowski, Jakub Michaliszyn

Institute of Computer Science,
University Of Wrocław,
ul. Joliot-Curie 15, 50-383 Wrocław, Poland

Abstract. It is well known, and easy to see, that not each nondeterministic Büchi automaton on infinite words can be simulated by a nondeterministic co-Büchi automaton. We show that in the cases when such a simulation is possible, the number of states needed for it can grow nonlinearly. More precisely, we show a nondeterministic Büchi automaton with n states, which can be simulated by a nondeterministic co-Büchi automaton, but cannot be simulated by any nondeterministic co-Büchi automaton with less than $c * n^{7/6}$ states for some constant c . This improves on the best previously known lower bound of $3n/2$.

O problemie istnienia języków zupełnych w klasach semantycznych

Zenon Sadowski
Instytut Matematyki, Uniwersytet w Białymstoku

Maszyna Turinga będąca modelem obliczeń semantycznej klasy złożoności obliczeniowej musi spełniać specyficzny dla tej klasy warunek zwany "promise".

W pracy badamy następujące problemy otwarte:

Q1: Czy istnieje optymalny system dowodowy dla danego języka L ?

Q2: Czy dana semantyczna klasa złożoności C posiada język zupełny ?

W przypadku konkretnych języków (takich jak $TAUT$ czy SAT) i konkretnych semantycznych klas złożoności (takich jak $\mathbf{NP} \cap \text{co-}\mathbf{NP}$, $\mathbf{UP}$, rozłączne $\mathbf{NP}$ -pary) problemy te były intensywnie badane. Nasze podejście polega na wyrażeniu warunku "promise" w języku L i używaniu systemu dowodowego dla L do weryfikowania czy "promise" jest spełniony. Przedstawiamy nową charakteryzację problemu Q2:

Semantyczna klasa złożoności C posiada problem zupełny wtedy i tylko wtedy, gdy język L , w którym warunek "promise" jest wyrażalny, posiada system dowodowy z krótkimi dowodami dla stwierdzeń mówiących, że maszyna Turinga spełnia "promise".

Literatura

- [1] Beyersdorff O., Sadowski Z., Characterizing the Existence of Optimal Proof Systems and Complete Sets for Promise Classes, Proc. 4th Computer Science Symposium in Russia (CSR 2009), Springer LNCS.

Problemy dotyczące trwałości w rozszerzeniach p/t-sieci*

Kamila Barylska

Uniwersytet Mikołaja Kopernika w Toruniu

Wydział Matematyki i Informatyki

khama@mat.uni.torun.pl

Klasyczne pojęcie trwałości (ang. *persistence*) pochodzące z lat 60-tych, będące uogólnieniem pojęcia bezkonfliktowości, jest do dziś intensywnie badane. Oparte jest ono o zasadę „nie odpychaj umożliwionego” - ten typ trwałości oznaczmy jako **e/e-P** (*enabled/enabled-Persistence*).

Krok MaM' jest **e/e-trwały** wtedy i tylko wtedy, gdy każda (różna od a) akcja umożliwiona w stanie M jest nadal umożliwiona w stanie M' .

Uogólnieniami klasycznego pojęcia trwałości są : **l/l-P** (*live/live-Persistence*) i **e/l-P** (*enables/live-Persistence*), oparte o zasadę „nie zabijaj”. Pierwsze, silniejsze – „nie zabijaj nikogo” i drugie, słabsze – „nie zabijaj umożliwionego”.

Krok MaM' nazywać będziemy **l/l-trwałym** wtedy i tylko wtedy, gdy każda (różna od a) akcja żywa w stanie M pozostaje żywa w stanie M' , natomiast **e/l-trwałym** wtedy i tylko wtedy, gdy każda (różna od a) akcja umożliwiona w stanie M jest żywa w stanie M' .

System nazywać będziemy e/e-trwałym (l/l-trwałym, e/l-trwałym) gdy każdy krok tego systemu będzie, odpowiednio, krokiem e/e-trwałym (l/l-trwałym e/l-trwałym, odpowiednio). Odpowiednie klasy systemów trwałych będziemy oznaczać odpowiednio: $P_{e/e}$, $P_{l/l}$ i $P_{e/l}$.

Wszystkie trzy problemy decyzyjne „Czy dana sieć Petriego jest x/y-trwała?” (gdzie $x/y = e/e, e/l$ bądź l/l) są rozstrzygalne w klasie markowanych sieci Petriego (*place/transition nets*).

System ma Własność Dużego Diamentu, jeżeli $MuM' \wedge MvM'' \wedge u \approx v$ (równoważność Parikh), to $M'=M''$, zaś Własność Diamentu, jeżeli $MabM' \wedge MbaM''$, to $M'=M''$.

W klasycznych p/t-sieciach trzy typy trwałości tworzą rosnącą hierarchię: $P_{e/e} \subset P_{l/l} \subset P_{e/l}$.

Pokażemy, że hierarchia trwałości jest prawdziwa w systemach „diamentowych”, natomiast załamuje się w rozszerzeniach p/t-sieci, które nie mają własności diamentu (np. w sieciach czyszczących lub podwajających).

*Badania wspierane przez Ministerstwo Nauki i Szkolnictwa Wyższego, grant N N206 258035

Unikanie wzorców w tekstach w dwóch wymiarach

Radosław Głowiński

Wydział Matematyki i Informatyki

Uniwersytet Mikołaja Kopernika

`glowir@mat.umk.pl`

Problem unikania wzorca w tekście jest szeroko rozpatrywanym w literaturze zagadnieniem. Poprzez tekst rozumiemy słowo nad skończonym alfabetem A , natomiast wzorzec to słowo nad alfabetem E rozłącznym z A . Dla ustalonego wzorca $p \in E^*$ oraz słowa $w \in A^*$, jeśli istnieje morfizm $h : E \rightarrow A^+$, tzn. $h(p)$ jest podslowem w , to mówimy, że słowo w zawiera wzorzec p . Istotnym pytaniem jest: dla jak małych alfabetów potrafimy stworzyć nieskończone słowo, w którym nie pojawia się dany wzorzec? Istnieje algorytm rozstrzygający czy dany wzorzec jest omijalny nad pewnym alfabetem (algorytm Zimina [1]). Aktualne rozważania są prowadzone dla tekstów i wzorców jednowymiarowych. Wprowadzamy pojęcie wzorca dwuwymiarowego oraz definiujemy sposób występowania takiego wzorca w tekście dwuwymiarowym. Przy szukaniu słów dwuwymiarowych unikających wzorca korzystamy z wyników dotyczących wzorców jednowymiarowych.

Literatura

- [1] Lotahire M.: *Algebraic Combinatorics on Words*
Encyclopedia of Mathematics and its Applications, vol. 90
Cambridge University Press, Cambridge 2002

Generowanie drzew binarnych w naturalnym porządku

Sebastian Smoczyński

Wydział Matematyki i Informatyki

Uniwersytet Mikołaja Kopernika

smyczek@mat.umk.pl

Drzewa jako hierarchiczne struktury danych stosowane są praktycznie w każdej dziedzinie informatyki i stanowią interesujący oraz wdzięczny temat badań. Podczas badania właściwości drzew binarnych jednym z naturalnie pojawiających się zagadnień jest problem efektywnego generowania wszystkich (nieetykietowanych) drzew binarnych o ustalonej liczbie wierzchołków. W roku 2006 Donald Knuth wydał cały zeszyt [1] poświęcony w całości temu zagadnieniu.

Zazwyczaj algorytmy generujące nie operują bezpośrednio na strukturze drzewa, lecz na pewnej jego reprezentacji. Najczęściej kodują drzewo ciągiem liczb i generują wszystkie prawidłowe ciągi dla danej reprezentacji w porządku leksykograficznym.

Na zbiorze wszystkich drzew binarnych $\mathcal{T}$ możemy zadać *naturalny porządek* uwzględniający rozmiar drzew [2] w taki sposób, aby drzewa o mniejszej ilości wierzchołków poprzedzały drzewa o większym rozmiarze. Rekurencyjnie definiujemy: $T_1 \prec T_2$ dla $T_1, T_2 \in \mathcal{T}$ jeżeli

1. $c(T_1) < c(T_2)$, lub
2. $c(T_1) = c(T_2)$ i $l(T_1) \prec l(T_2)$, lub
3. $c(T_1) = c(T_2)$ i $l(T_1) = l(T_2)$ i $r(T_1) \prec r(T_2)$,

gdzie $c(T)$ oznacza rozmiar drzewa T (ilość wierzchołków), $l(T)$ jest lewym poddrzewem drzewa T , zaś $r(T)$ prawym poddrzewem T . Oczywiście porządek ten jest dobrze zdefiniowany również w zbiorze $\mathcal{T}_n$ wszystkich drzew binarnych o ustalonej liczbie n wierzchołków.

Etykietując wierzchołki drzewa kolejnymi liczbami naturalnymi przy przechodzeniu drzewa w porządku inorder i wypisując te etykiety przy ponownym przejściu tym razem w porządku preorder otrzymuje się permutację będącą *reprezentacją preorder* danego drzewa. Okazuje się, że porządek leksykograficzny na permutacjach preorder zachowuje naturalny porządek reprezentowanych drzew binarnych.

Liniowy, rekurencyjny algorytm generujący wszystkie drzewa binarne o ustalonej liczbie wierzchołków bazujący na reprezentacji preorder drzew został opisany w pracy [3].

W referacie przedstawiony zostanie nowy, liniowy i nierekurencyjny algorytm generujący wszystkie drzewa w naturalnym porządku.

Literatura

- [1] Knuth, D.E.: The Art of Computer Programming, Volume 4: Generating All Trees. Addison Wesley (2006)
- [2] Knott G.D.: A numbering system for binary trees. Commun. ACM **20**(2) (1977) 113–115
- [3] Solomon M.H., Finkel R.A.: A note on enumerating binary trees. J. ACM **27**(1) (1980) 3–5

Otrzymywanie klucza kryptograficznego z pewnych prawie separowalnych stanów kwantowych o związanym splątaniu

Łukasz Pankowski ^{*†‡}

Kryptografia kwantowa pozwala na uzyskanie bezpieczeństwa opartego na fizycznej niemożliwości podsłuchiwania — odczyt przesyłanego sygnału przez podsłuchiacza jest nierozdzielnie związany z wprowadzeniem szumu do przesyłanego sygnału, zatem Alicja i Bob mogą wyznaczyć poziom szumu w kanale i sprawdzić, czy jest on wystarczająco słaby by zagwarantować im bezpieczeństwo (nawet jeśli szum w całości byłby wynikiem działania podsłuchiacza). Wyodrębniamy dwa typy protokołów kwantowych: *przygotuj i zmierz* (np. słynny protokół BB84 [1]) oraz protokoły oparte na przesyłaniu „połówek” stanów splątanych (np. równie słynny protokół Ekerta [2]). Dowody bezpieczeństwa protokołów *przygotuj i zmierz* przeprowadzono pokazując ich równoważność z destylacją stanów maksymalnie splątanych (tzw. *singletów*). Doprowadziło to do przeświadczenia, że bezpieczeństwo kwantowej kryptografii jest równoważne destylacji singletów.

Takie przeświadczenie sugerowało, że ze stanów o *związanym splątaniu* [3] (tzn. stanów splątanych, z których nie można wydestylować singletów) nie da się otrzymać bezpiecznego klucza kryptograficznego. Wbrew powszechnemu przeświadczeniu udało się uzyskać klucz ze stanów o *związanym splątaniu* [4] nawet dla stanów nisko wymiarowych [5].

Stany nisko wymiarowe dla których w [5] uzyskano bezpieczny klucz znajdują się na granicy stanów niedestylowalnych¹, a istnienie stanów o tej własności wewnątrz klasy stanów niedestylowalnych pokazano za pomocą argumentu o ciągłości klucza, bez wskazania analitycznej formuły takich stanów.

Podczas referatu zaprezentuję klasę stanów o związa-

nym splątaniu z których można otrzymać klucz kryptograficzny. Stany z prezentowanej klasy podchodzą dowolnie blisko klasy stanów separowalnych [6], z których nie można uzyskać klucza kryptograficznego. Zaprezentuję również prosty warunek wystarczający destylowalności klucza kryptograficznego [6], za pomocą którego można wykazać, że z prezentowanej klasy stanów można destylować klucz kryptograficzny.

Literatura

- [1] C. H. Bennett and G. Brassard, “Quantum cryptography: Public key distribution and coin tossing,” in *Proceedings of the IEEE International Conference on Computers, Systems and Signal Processing*. Bangalore, India, December 1984: IEEE Computer Society Press, New York, 1984, pp. 175–179.
- [2] A. K. Ekert, “Quantum cryptography based on Bell’s theorem,” *Phys. Rev. Lett.*, vol. 67, pp. 661–663, 1991.
- [3] M. Horodecki, P. Horodecki, and R. Horodecki, “Mixed-state entanglement and distillation: Is there a “bound” entanglement in nature?” *Phys. Rev. Lett.*, vol. 80, pp. 5239–5242, 1998, quant-ph/9801069.
- [4] K. Horodecki, M. Horodecki, P. Horodecki, and J. Oppenheim, “Secure key from bound entanglement,” *Phys. Rev. Lett.*, vol. 94, p. 160502, 2005, quant-ph/0309110.
- [5] K. Horodecki, Ł. Pankowski, M. Horodecki, and P. Horodecki, “Low dimensional bound entanglement with one-way distillable cryptographic key,” *IEEE Trans. Inf. Theory*, vol. 54, pp. 2621–2625, 2008, arXiv:quant-ph/0506203.
- [6] Ł. Pankowski and M. Horodecki, in preparation.

^{*}Instytut Informatyki, Uniwersytet Gdański, Gdańsk

[†]Instytut Fizyki Teoretycznej i Astrofizyki, Uniwersytet Gdański, Gdańsk

[‡]Krajowe Centrum Kwantowej Informatyki w Gdańsku, Sopot

¹Właściwie to na granicy tzw. stanów PPT. Pytanie czy wszystkie stany poza PPT są destylowalne wciąż pozostaje otwarte.

O rodzajach E -unifikacji

Jan Otop

25 kwietnia 2009

Niech E będzie teorią równościową nad sygnaturą Σ_1 . E -unifikacja to problem znalezienia dla zbioru równań na termach Γ nad sygnaturą Σ_2 , takiego podstawienia Θ , że dla każdego $s = t \in \Gamma$ zachodzi $\Theta(s) =_E \Theta(t)$, gdzie $s =_E t$ oznacza $E \vdash s = t$.

Rozróżniamy trzy rodzaje E -unifikacji ze względu na ograniczenia na sygnaturę Σ_2 :

- elementarna E -unifikacja, gdy $\Sigma_1 = \Sigma_2$
- E -unifikacja ze stałymi, gdy $\Sigma_2 = \Sigma_1 \cup \mathcal{C}$ gdzie $\mathcal{C}$ jest zbiorem stałych, które nazywamy stałymi wolnymi ponieważ nie występują w E
- ogólna E -unifikacja, gdy nie ma ograniczeń na Σ_2

Teoria AC (łączność i przemienność) wydaje się dość prosta, do jej opisu wystarczą dwa równania na każdy operator który ma być łączny i przemienny. Jednakże AC -unifikacja, używana w systemie dowodzenia twierdzeń EQP, była badana przez ponad 10 lat zanim udowodniono że podany algorytm jest poprawny i podano optymalny algorytm generujący zupełny zbiór unifikatorów. Co zaskakujące, pojedyncze równanie może mieć podwójnie wykładniczo wiele rozwiązań względem swojej długości. Świadczy to o tym, że problem E -unifikacji nawet dla dość prostych teorii jest skomplikowany oraz konstrukcja algorytmu jest niebanalna. Oczywiście, bardzo łatwo wskazać teorię dla których E -unifikacja jest nierozstrzygalna. Będziemy się jednak zajmowali teoriami z rozstrzygalną E -unifikacją.

Często, im więcej jest aksjomatów w teorii E , tym trudniej skonstruować algorytm rozstrzygający E -unifikację. Pojawia się pytanie, czy możliwe jest skonstruowanie algorytmu unifikacji dla sumy teorii $E_1 \cup E_2$ z dwóch algorytmów E_1 - i E_2 -unifikacji? Pytanie to cieszyło się dużą popularnością w latach 90 i zostały podane algorytmy dla przypadku, gdy sygnatury E_1 i E_2 są rozłączne oraz spełnione są dodatkowe założenia, mianowicie E_1 - i E_2 -unifikacja z ograniczeniami liniowymi musi być rozstrzygalna.

Definicja 1. Niech E będzie teorią równościową nad sygnaturą Σ oraz $\mathcal{C}$ będzie przeliczalnym zbiorem stałych nie występujących w Σ . Instancją problemu E -unifikacji z ograniczeniami liniowymi jest zbiór równań na termach Γ nad sygnaturą $\Sigma \cup \mathcal{C}$ oraz liniowy porządek $<$ na zbiorze stałych wolnych (stałych z $\mathcal{C}$) i zmiennych występujących w Γ . Definiujemy $V_c = \{X : X \text{ jest zmienną oraz } X < c\}$ dla każdej stałej wolnej występującej w Γ . Rozwiązaniem problemu jest podstawienie Θ , takie że

- Θ jest rozwiązaniem każdego równania z Γ
- jeśli $X \in V_c$ to $\Theta(X)$ nie zawiera c

Pozostało pytanie, czy dodatkowe założenia są konieczne? Problem, czy istnieje teoria E , taka że E -unifikacja ze stałymi jest rozstrzygalna, natomiast E -unifikacja z ograniczeniami liniowymi jest nierozstrzygalna znalazł się na liście RTA jako problem numer 66. Na liście RTA znajduje się też powiązany problem, mianowicie, czy istnieje teoria E , taka że E -unifikacja z ograniczeniami liniowymi jest rozstrzygalna a E -unifikacja z ogólnymi ograniczeniami (tj. zbiory V_c są dowolnymi podzbiórmi zbioru zmiennych, niekoniecznie definiowane przy pomocy porządku liniowego) jest nierozstrzygalna.