

4. Контекстно-свободни граматика и езици. Стекови автомати.

1. Д Контекстно-свободна граматика G наричаме четворката (V, Σ, R, S) , където:

- V е крайна азбука,
- Σ (множесиво от терминали) е подмножесиво на V ,
- R (множесиво от правила) е крайно подмножесиво на $(V - \Sigma) \times V^*$, и
- S (начален символ) е елемент на $V - \Sigma$.

Елементите на $V - \Sigma$ се наричат нестерминали. За всяко $A \in V - \Sigma$ и $u \in V^*$ казваме, че $A \rightarrow_G u$ вита, когато $(A, u) \in R$. За всеки две думи $u, v \in V^*$, пишем $u \Rightarrow_G v$ т.е.т.х. има думи $x, y \in V^*$ и $A \in V - \Sigma$ такива че $u = xAy$, $v = xv'y$ и $A \rightarrow_G v'$. Релацията $\Rightarrow_G^*$ е рефлексивна, транзитивна затвореност на $\Rightarrow_G$. И накрая, $L(G)$, езикът генериран от G е $\{w \in \Sigma^* : S \Rightarrow_G^* w\}$; казваме още, че G генерира всяка дума от $L(G)$. Едни език L се нарича контекстно-свободен ако $L = L(G)$, за някоя КСГ G .

Когато граматиката се подразбира пишем $A \rightarrow w$ и $u \Rightarrow v$, вместо $A \rightarrow_G w$ и $u \Rightarrow_G v$.

Наричаме всяка последователност от вита:

$$w_0 \Rightarrow_G w_1 \Rightarrow_G w_2 \Rightarrow_G \dots \Rightarrow_G w_n$$

извод в G на w_n от w_0 . Тук $w_0, \dots, w_n$ са произволни думи от V^* да n наричаме дължина на извода, $n \in \mathbb{N}$ и $n \geq 0$. Казваме още, че изводът има n стъпки.

Пример Следната граматика генерира всички изрази с правилно вложени скоби: всяка лява скоба трябва да има съответстваща дясна, и всяка затваряща гръба да има уникнална съответстваща отиваща гръба. Още по-лесно, назовем мемлу всяка такава двойка скоби трябва да отиващи на същото условие. Нека $G = (V, \Sigma, R, S)$,

$$\begin{aligned} V &= \{S, (,)\}, \\ \Sigma &= \{(,)\}, \\ R &= \{S \rightarrow e, \\ &\quad S \rightarrow SS, \\ &\quad S \rightarrow (S)\}. \end{aligned}$$

Два извода от тази граматика са:

$$S \Rightarrow SS \Rightarrow S(S) \Rightarrow S((S)) \Rightarrow S(()) \Rightarrow ()()$$

$$S \Rightarrow SS \Rightarrow (S)S \Rightarrow ()S \Rightarrow ()(S) \Rightarrow ()()$$

Следователно една дума може да има няколко различни извода в една граматика. $L(G)$ е пример за език, който е контекстно-свободен, но не е регулярен.

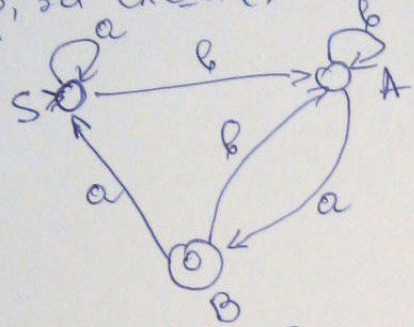
Очевидно има примери за езици, които са контекстно-свободни, но не са регулярни. Всички регулярни езици обаче са контекстно-свободни. Заста ще покажем това чрез директно построение:

Нека разгледаме регулярния език, който се разпознава от крайния детерминистичен автомат $M = (K, \Sigma, \delta, s, F)$. Същият език се генерира от граматиката $G(M) = (V, \Sigma, R, S)$, където $V = K \cup \Sigma$, $S = s$ и R се състои от тези правила:

$$R = \{q \rightarrow aq : \delta(q, a) = q\} \cup \{q \rightarrow \epsilon : q \in F\}$$

Тоест: детерминистите са състоянията на автомата; а за всеки преход от q към r при вход a имаме в R правило $q \rightarrow ar$.

Например, за следния автомат:



ще построим граматиката:

$$S \rightarrow aS, S \rightarrow bA, A \rightarrow aB, A \rightarrow bA, B \rightarrow aS, B \rightarrow bA, B \rightarrow \epsilon.$$

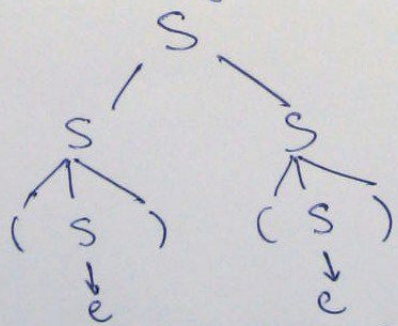
② Дървета на извода / Дървета за синтактичен анализ

Нека G е КСГ. Както показвахме, думата $w \in L(G)$ може да има много изводи в G . Например в разглежданата по-горе граматика за правилно вложени скоби, думата $()()$ може да бъде получена от S чрез поне два различни извода:

$$S \Rightarrow SS \Rightarrow (S)S \Rightarrow ()S \Rightarrow ()(S) \Rightarrow ()() \quad , \text{ и}$$

$$S \Rightarrow SS \Rightarrow S(S) \Rightarrow (S)(S) \Rightarrow ()(S) \Rightarrow ()()$$

Ако разгледаме тези два извода ще забележим, че те са еднакви-разгледаме се единично по реда на прилагане на правилата. Интуитивно и двете могат да бъдат означени графично така:

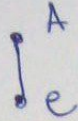


Диаграми от горния тип наричаме дървета на извода. Тогатите наричаме върхове; всеки връх е означен със символ от V . Най-горният връх се нарича корен, а най-долният-лист. Всички лист са означени или с терминални символи, или с празната дума ϵ . Ако следим местата на дървото ще получим дума, която се нарича **резултат** на дървото на извода.

Или казано по-формално, за произволна КСГ $G = (V, \Sigma, R, S)$ дърво на извода, корен, листо, ~~yield~~ ^{резултат}, така:

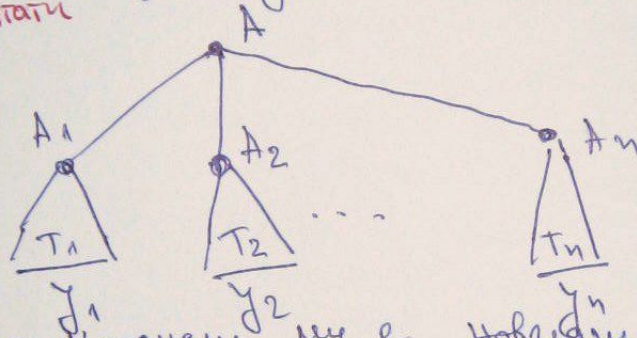
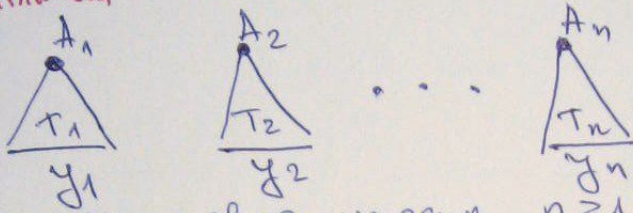
(1) Това е дърво на извода за всяко $a \in \Sigma$. Единственият връх е и корен, и листо. ~~yield~~ ^{резултат} е "a".

(2) Ако $A \rightarrow e$ е правило в R , то



е дърво на извод. Коренът му е означен с A , единственият му листо е e , ~~yield~~ ^{резултат} е e .

(3) Ако $A_1, \dots, A_n$ са дървета на извод, където $n \geq 1$, с корени обозначени с $A_1, \dots, A_n$ и ~~yield~~ ^{резултат} $y_1, \dots, y_n$ и $A \rightarrow A_1 \dots A_n$ е правило от R , то



е дърво на извода. Коренът му е новият връх A , листата му са листата на дърветата $A_1, \dots, A_n$, ~~yield~~ ^{резултат} е $y_1 \dots y_n$.

(4) Има други дървета на извода.

Изменливо, дърветата на извода представят изводите по такъв начин, че различни дървета се наредят на прилагане на правилата, се симулират. Казано иначе, дърветата на извода дефинират класове на еквивалентност на изводите. По формално: Нека $G = (V, \Sigma, R, S)$ е КСГ, $D = x_1 \Rightarrow x_2 \Rightarrow \dots \Rightarrow x_n$ и $D' = x'_1 \Rightarrow \dots \Rightarrow x'_n$ са два извода в G , където $x_i, x'_i \in V^*$ за $i = 1, \dots, n$, $x_1, x'_1 \in V - \Sigma$ и $x_n, x'_n \in \Sigma^*$. Тоеест, и двата са изводи на терминални символи от един нетерминал. Казваме, че D предхожда D' (записваме $D \prec D'$), ако $n > 2$ и съществува $k \in \mathbb{N}$, $1 < k < n$, такова че:

- (1) За всяко $i \neq k$ $x_i = x'_i$
- (2) $x_{k-1} = x'_{k-1} = uAvBw$, където $u, v, w \in V^*$ и $A, B \in V - \Sigma$
- (3) $x_k = uywBw$, където $A \rightarrow y \in R$
- (4) $x'_k = uAvzw$, където $B \rightarrow z \in R$
- (5) $x_{k+1} = x'_{k+1} = uyvzw$

С други думи, двата извода са идентични, с изключение на две последователни шизми, през които едни и същи две неутронизирани са заменени с едни и същи думи, но в променен ред. Изводи, в които по-левият неутронизирател е

Пример Да разгледаме следните 3 извода Φ_1, Φ_2 и Φ_3 може без примера

$$D_1 = S \Rightarrow SS \Rightarrow (S)S \Rightarrow ((S))S \Rightarrow (((S)))S \Rightarrow (((S)))((S)) \Rightarrow (((S)))((S))((S))$$

$$D_2 = S \Rightarrow SS \Rightarrow (S)S \Rightarrow ((S))S \Rightarrow ((S))(S) \Rightarrow (())(S) \Rightarrow (())()$$

$$\Phi_3 = S \Rightarrow SS \Rightarrow (S)S \Rightarrow ((S))S \Rightarrow ((S))(S) \Rightarrow ((S))(\cdot) = \chi((\cdot))(\cdot)$$

Имаме, че $D_1 \prec D_2$ и $D_2 \prec D_3$, но $D_1 \prec D_3$ НЕ е извлечено, тъй като D_1 и D_3 се размишават в повече от една машинна дума. И трите имат едно и също дърво на извода.

Казваме, че два извода D и D' са подобни, ако двойката (D, D') принадлежи на рефлексивното, симетрично, транзитивно затворено отношение $\sim$ (свързано с отношението на еквивалентност).

Всички такива клас на еквивалентност (т.е. всяко дърво на извода), съдържа извод, който е максимален спрямо $\prec$ (не го предхожда никой друг извод). Нарича се най-голям извод и се поглежда като от корена. А последователно заменяме най-големия некрещилен в думата според правилото, означено в дървото. Подобно се достига най-десен извод - извод, който не предхожда никой друг извод и се поглежда чрез последователно разширяване на най-десния номериран. Всяко дърво на извода съдържа точно един най-голям и точно един най-десен извод.

за да обобщим, ще формулираме следната теорема:

Т] Нема $G = (V, \Sigma, R, S)$ с хсг и $A \in V - \Sigma$ и $w \in \Sigma^*$. Следните твърдения са еквивалентни:

(a) $A \supset B \wedge C$

(6) ~~Che~~ ~~ne~~ ~~ca~~ ~~zapo~~ ~~ro~~ na izboga s koret A u ~~izida w.~~

(c) ~~Тяжелый~~ Най-голям излог $A \xrightarrow{Z} w$

(d) Вещество H_2O - жидк. изобр. $A \xrightarrow{P} W$

Снежови Авшомани

Не всеки контекстно-свободен език може да бъде разпознат от краен автомат, понеже не всеки контекстно-свободен език е регулярен. За да получим автомати, които да разпознават контекстно-свободните езици трябва да добавим нови възможности към крайните автомати. Оказва се, че ключовата особеност на тези по-сложни автомати е добавянето на допълнителна памет, под формата на стек.

DI Ще дефинираме стехов автомат като наредената шесторка

$M = (K, \Sigma, \Gamma, \Delta, s, F)$, където:

- K е крайно множество от състояния
- Σ е ^{крайна} азбука (входни символи)
- Γ е ^{крайна} азбука (символи на стека)
- $s \in K$ е началното състояние
- $F \subseteq K$ е множество от крайни състояния, и
- Δ , релация на прехода, е крайно подмножество на:
 $(K \times (\Sigma \cup \{ \epsilon \}) \times \Gamma^*) \times (K \times \Gamma^*)$.

Изтълкуваме, ако $((p, a, \beta), (q, \gamma)) \in \Delta$, това M , винаги когато е в състояние p с β на върха на стека, може да прочете a от входа (или нищо, ако $a = \epsilon$), да замени β с γ на върха на стека, и да премине към състояние q . Така двойката $((p, a, \beta), (q, \gamma))$ се нарича преход в M . Тъй като в даден момент може да има няколко възможни прехода, стеховите автомати по принцип са недетерминистични.

Дефинираме две операции: $push$ (добавяне на ~~символ~~ на върха на стека), и pop (премахване на символ от върха на стека).

Например преходът $((p, u, \epsilon), (q, a))$ демонстрира $push$, а $((p, u, a), (q, \epsilon))$ - pop .

Конфигурация на стехов автомат се дефинира като елемент на $K \times \Sigma^* \times \Gamma^*$: първият компонент е текущото състояние, вторият е частта от лентата, която още не е била прочетена, а третият е съдържанието на стека, прочетено отгоре надолу.

Ако (p, x, α) и (q, y, β) са две конфигурации на M , казваме се (p, x, α) води след една стъпка до (q, y, β) (бележим $(p, x, \alpha) \vdash_M (q, y, \beta)$) ако съществува преход $((p, a, \beta), (q, \gamma)) \in \Delta$, такова че $x = a\gamma$, $\alpha = \delta\gamma$ и $\beta = \gamma\gamma'$ за някое $\gamma' \in \Gamma^*$. Бележим рефлексивно и

транзитивно затваряне на $\vdash_M$ с $\vdash_M^*$. Казваме се M разпознава думата $w \in \Sigma^*$ т.с.т.к. $(s, w, \epsilon) \vdash_M^* (p, \epsilon, \epsilon)$ за някое $p \in F$. Казано иначе M разпознава w т.с.т.к. $\exists$ последователност от конфигурации

$C_0, C_1, \dots, C_n$ ($n \geq 0$) такава че $C_0 \vdash_M C_1 \vdash_M \dots \vdash_M C_n$, $C_0 = (s, w, \epsilon)$, $C_n = (p, \epsilon, \epsilon)$ за някое $p \in F$. Всяка последователност от конфигурации

$C_0, C_1, \dots, C_n$, такава че $C_i \vdash_M C_{i+1}$, $i = 0, \dots, n-1$ се нарича изчисление от M ; дължината му е n , или има n стъпки. Езикът, разпознат от M , бележим с $L(M)$, е множество от всички думи, разпознавани от M .

Когато не може да настъпи отбавяне пишем $\vdash$ и $\vdash^*$ вместо $\vdash_M$ и $\vdash_M^*$.

5 II Класът на езиците, разпознавани от стехови автомати, съвпада с класа на контекстно свободните езици.

Доказателство: Ще разделим доказателството на две части

=>)

Лема | Всеки контекстно-свободен език се разпознава от някой стенов автомат.

Доказателство:

Нека $G = (V, \Sigma, R, S)$ е КСГ; трябва да построим стенов автомат M , така че $L(M) = L(G)$. Ще построим автомат M с точно две състояния, p и q , който свои входи в състояние q след първия си приход. M използва V , множеството от терминиали и не-терминиали, за азбуката на стена. Нека $M = (\{p, q\}, \Sigma, \Delta, p, \{q\})$ където Δ съдържа следните преходи:

- (1) $((p, e, e), (q, S))$
 - (2) $((q, e, A), (q, x))$ за всяко правило $A \rightarrow x$ в R
 - (3) $((q, a, a), (q, e))$ за всяко $a \in \Sigma$ (преход тип 1)
- Стеновият автомат първо сваля S на върха на стена, и след това последователно или

- заменя най-горния символ A в стена, ако е не-терминал, с дясната страна x на някое правило $A \rightarrow x$ (преход тип 2), или

- премахва символа на върха на стена ако е терминал и съвпада със следващия символ на лентата (преход тип 3).

По този начин M симулира най-ляв извод на входа.

За да докажем, че $L(M) = L(G)$, ще използваме следното твърдение:

Твърдение | Нека $w \in \Sigma^*$ и $\alpha \in (V - \Sigma)V^*V\{e\}$. Тогава $S \xRightarrow{L^*} w\alpha$ т.с.т.к.

$(q, w, S) \vdash_M^* (q, e, \alpha)$

Това твърдение ще ни помогне за доказване лемата, тъй като ако $a = e$ ще следва, че $S \xRightarrow{L^*} w$ т.с.т.к. $(q, e, S) \vdash_M^* (q, e, e)$, т.е. $w \in L(G)$ т.с.т.к. $w \in L(M)$.

$\Rightarrow$ Нека $S \xRightarrow{L^*} w\alpha$, където $w \in \Sigma^*$ и $\alpha \in (V - \Sigma)V^*V\{e\}$. Ще докажем по индукция, че по дължината на най-левия извод на w от S , че $(q, w, S) \vdash_M^* (q, e, \alpha)$.

1. Ако дължината на извода е 0, тогава $w = e$, $\alpha = S$ и наистина $(q, w, S) \vdash_M^* (q, e, \alpha)$

2. Допускаме, че ако $S \xRightarrow{L^*} w\alpha$ при извод с дължина n , или по-малко, $n \geq 0$, то $(q, w, S) \vdash_M^* (q, e, \alpha)$.

3. Нека $S = u_0 \xRightarrow{L} u_1 \xRightarrow{L} \dots \xRightarrow{L} u_n \xRightarrow{L} u_{n+1} = w\alpha$ да е най-ляв извод на $w\alpha$ от S . Нека A е най-леви не-терминал на u_n . Тогава $u_n = xAv$ и $u_{n+1} = x\gamma v$, където $x \in \Sigma^*$, $\gamma, v \in V^*$ и $A \rightarrow \gamma$ е правило в R . Тъй като съществува най-ляв извод с дължина n на $u_n = xAv$ от S , то по И.П. $(q, x, S) \vdash_M^* (q, e, Av)$ (2)

Той като $A \rightarrow \gamma$ е правило в R :

$$(q, e, \alpha\beta) \vdash_M (q, e, \gamma\beta) \quad (3)$$

след преход от тип 2.

Той като $n+1$ е цяло, но и $x\gamma\beta$, има дума $y \in \Sigma^*$, такава че $w = xy$ и $y\alpha = \gamma\beta$. Тогава може да предположим (2) и (3) така:

$$(q, w, S) \vdash_M^* (q, y, \gamma\beta) \quad (4)$$

Но $y\alpha = \gamma\beta \Rightarrow$

$$(q, y, \gamma\beta) \vdash_M^* (q, e, \alpha) \quad (5)$$

след последователност от $|y|$ на двой прехода от тип 3. Обединяването на (4) и (5) прилага индукцията.

$\Leftarrow$) Сега да допуснем, че $(q, w, S) \vdash_M^* (q, e, \alpha)$ с $w \in \Sigma^*$ и $\alpha \in (V - \Sigma) \cup \{e\}$. Ще покажем, че $S \stackrel{*}{\Rightarrow} w\alpha$. Отново ще използваме индукция, този път по броя преходи от тип 2.

1. Той като първия преход винаги е от тип 2, ако $(q, w, S) \vdash_M^* (q, e, \alpha)$ без преходи от тип 2, то $w = e$ и $\alpha = \beta \Rightarrow$ верно

2. Ако $(q, w, S) \vdash_M^* (q, e, \alpha)$ след n стъпки от тип 2 (или по-малко), $n \geq 0$, то $S \stackrel{*}{\Rightarrow} w\alpha$

3. Да предположим че $(q, w, S) \vdash_M^* (q, e, \alpha)$ с $n+1$ прехода от тип 2, и да разгледаме последния такъв преход:

$$(q, w, S) \vdash_M^* (q, y, \alpha\beta) \vdash_M (q, y, \gamma\beta) \vdash_M^* (q, e, \alpha)$$

където $w = xy$ за някои $x, y \in \Sigma^*$ и $A \rightarrow \gamma$ е правило от граматиката. Според И.Х. имаме че $S \stackrel{*}{\Rightarrow} x\alpha\beta$ и $\Rightarrow S \stackrel{*}{\Rightarrow} x\gamma\beta$. Но поеме $(q, y, \gamma\beta) \vdash_M^* (q, e, \alpha)$, евентуално след преход от тип 3, следва че $y\alpha = \gamma\beta$ и следователно $S \stackrel{*}{\Rightarrow} x\gamma\alpha = w\alpha$. Това лемата е доказана, както и първата половина на теоремата.

Нека припомним към втората половина на теоремата:

Лема $\Leftarrow$) Ако един език се разпознава от стекъв автомат, то той е контекстно свободен.

Доказателство:

Идеята ще е да ограничим някак стекъв автомат, който разпознава. Ще наречем стекъв автомат прост ако:

Винаги когато $((q, \alpha, \beta), (p, \gamma))$ е преход в стекъв автомат и q не е началното състояние, то $\beta \in \Gamma$ и $|\gamma| \leq 2$.

с други думи виван се разпознава само символа на върха на стека и са заменят или с e , или с 1 или с 2 стенови символа.

Твърдим, че ^{ако} един език се разпознава от обикновен стенов автомат, то той се разпознава и от прост стенов автомат. За целта нека $M = (K, \Sigma, \Gamma, \Delta, S, F)$ е произволен стенов автомат. Ще построим $M' = (K', \Sigma, \Gamma \cup \{Z\}, \Delta', S', F')$ - прост, който също разпознава $L(M)$.
 Тук S' и F' са нови състояния, не принадлежащи в K , а Z е нов стенов символ, елемент на стека (сощо принадлежащ в Γ).
 Запознаваме като към Δ^* добавяме преходите:

- 1) $((S', e, e), (S, Z))$
- 2) $((f, e, Z), (f', e))$ за $\forall f \in F$.

Така Δ' съдържа новите 2 прехода, както и всички преходи от Δ . Следващата ни стъпка ще е да заменим в Δ' всички преходи, които не удовлетворяват условия за простота, с еквивалентни преходи, които го удовлетворяват. Това ще стане на 3 етапа:

- Преходи с $|p| \geq 2$
- Преходи с $|p| > 2$ (без да добавяме преходи с $|p| \geq 2$)
- Преходи с $p = e$ (без да нарушаване преходите 2).

1) За всички преходи от типа $((q, u, p), (r, \delta)) \in \Delta'$ и $|p| \geq 1$, $B = B_1 \dots B_n$, $n > 1$, добавяме към Δ' :

$$((q, e, B_1), (q_{B_1}, e)),$$

$$((q_{B_1}, e, B_2), (q_{B_1 B_2}, e)),$$

$$((q_{B_1 \dots B_{n-1}}, e, B_n), (r, \delta));$$

където за $i=1 \dots n-1$, $q_{B_1 B_2 \dots B_i}$ е ново състояние със значение "състояние q след премахване на символите $B_1 \dots B_i$ от стека".
 Очевидно новият автомат е еквивалентен на оригиналния.

2) По подобен начин добавяме преходите $((q, a, p), (r, \delta))$, $\delta = a \dots c$, $m \geq 2$ с

$$((q, a, p), (r_1, c_m))$$

$$((r_1, e, e), (r_2, c_{m-1}))$$

⋮

$$((r_{m-2}, e, e), (r_{m-1}, c_2))$$

$$((r_{m-1}, e, e), (r, c_1))$$

$r_1 \dots r_{m-1}$ са нови състояния. Нещо повече, са $\forall$ преходи $((q, a, p), (r, \delta)) \in \Delta'$ изпълняват $|p| \leq 1$. Ще го възстановим за $|p| \leq 2$.
 На следващата стъпка. Освен това, не добавяме преходи с $|p| \geq 1$.

3) Последно, нека разгледаме преходи от вида $((q, a, c), (p, b))$ с $q \neq s'$ - само те всъщност нарушават условието за простота. Ще заменим тези преходи с $((q, a, A), (p, \lambda A))$ за $\forall A \in \Gamma \cup \{e, z\}$. Сега, автоматът може да се премества без да поменя стека, то той може да се премества и ако последен последен символ и вестна го валиден.

Лесно се вижда, че така получаване прост автомат M' и $L(M) = L(M')$. За да докажем лемата ще покажем КСГ G , такава че $L(G) = L(M')$.

Нека $G = (V, \Sigma, R, S)$. V съдържа освен новия символ $\$$ и символите в Σ , нов символ $\langle q, A, r \rangle$ за $q, r \in K'$ и за всяко $A \in \Gamma \cup \{e, z\}$. За да се разбере ролята на термините $\langle q, A, r \rangle$ трябва да се има предвид, че G трябва да генерира всички думи, разпознавани от M' . $\Rightarrow$ термините на G представляват разпознаваеми части на входните думи, които се разпознават от M' . В частност, ако $A \in \Gamma$, тогава терминът $\langle q, A, r \rangle$ представлява която и да е част от входния низ, която може да бъде прочетена между момента, в който M' е в състояние q и A на върха на стека, и момента, в който M' премахва A от върха на стека и преминава в състояние r . Ако $A = e$, то $\langle q, e, r \rangle$ означава част от входната дума, която може да се прочете между състояние q и r , без да се прочега или променя.

Правилата в R са от 4 типа:

- (1) Правилото $\$ \rightarrow \langle s, z, f' \rangle$, където s е началното състояние на M , а f' - новото крайно състояние.
- (2) За всеки преход $((q, a, b), (r, c))$, където $q, r \in K'$, $a \in \Sigma \cup \{e, z\}$, $b, c \in \Gamma \cup \{e, z\}$ и за всяко $p \in K'$, добавяме правилото $\langle q, b, p \rangle \rightarrow a \langle r, c, p \rangle$.
- (3) За всеки преход $((q, a, b), (r, c_1, c_2))$, където $q, r \in K'$, $a \in \Sigma \cup \{e, z\}$, $b \in \Gamma \cup \{e, z\}$, и $c_1, c_2 \in \Gamma$ и за всяко $p, p' \in K'$, добавяме правилото $\langle q, b, p \rangle \rightarrow a \langle r, c_1, p' \rangle \langle p', c_2, p \rangle$.
- (4) За $\forall q \in K'$, правилото $\langle q, e, q \rangle \rightarrow e$.

Той като M' е прост, или тип 2 или тип 3 отговаря на всеки преход от M' .

Твърдение: За $\forall q, r \in K'$, $A \in \Gamma \cup \{e, z\}$ и $x \in \Sigma^*$,
 $\langle q, A, r \rangle \Rightarrow_G^* x$ т.с.т.к. $(q, x, A) \vdash_{M'} (r, e, e)$.

Нематом, а с това и теоремата, следват от това твърдение
 поемне $\langle S, e, f \rangle \Rightarrow_G^* x$ за няка $f \in F$ т.с.т. $(S, x, e) \vdash_{M'}^* (f, e, e)$
 т.е. $x \in L(G)$ т.с.т. $x \in L(M')$.

И двете посоки на твърдението могат да бъдат доказани
 по индукция, или по дължината на изводи на G ,
 или по дължината на изчисленията на M' .

6

Т Контекстно-свободните езици са затворени спрямо обединение,
 спелване и звезда на Клини.

Доказателство:

Нека $G_1 = (V_1, \Sigma_1, R_1, S_1)$ и $G_2 = (V_2, \Sigma_2, R_2, S_2)$ са две КСГ. БОО
 можем да смятаме, че $V_1 - \Sigma_1$ и $V_2 - \Sigma_2$ нямат общи елементи.

Обединение: Нека S е нов символ и нека $G = (V_1 \cup V_2 \cup \{S\}, \Sigma_1 \cup \Sigma_2, R, S)$,
 където $R = R_1 \cup R_2 \cup \{S \rightarrow S_1, S \rightarrow S_2\}$. Тогова твърдим, че
 $L(G) = L(G_1) \cup L(G_2)$. Поемне единствените правила, в които участва
 S , са $S \rightarrow S_1$ и $S \rightarrow S_2$, то $S \Rightarrow_G^* w$ т.с.т. x или $S_1 \Rightarrow_G^* w$, или
 $S_2 \Rightarrow_G^* w$; и т.к. G_1 и G_2 имат ~~не~~ множества от нетерминали без
 общи елементи, последното твърдение е еквивалентно на $w \in L(G_1) \cup L(G_2)$.

Спелване: Подхождаме по подобен начин: $L(G_1) L(G_2)$ се генерира
 от граматиката:

$$G = (V_1 \cup V_2 \cup \{S\}, \Sigma_1 \cup \Sigma_2, R_1 \cup R_2 \cup \{S \rightarrow S_1 S_2\}, S)$$

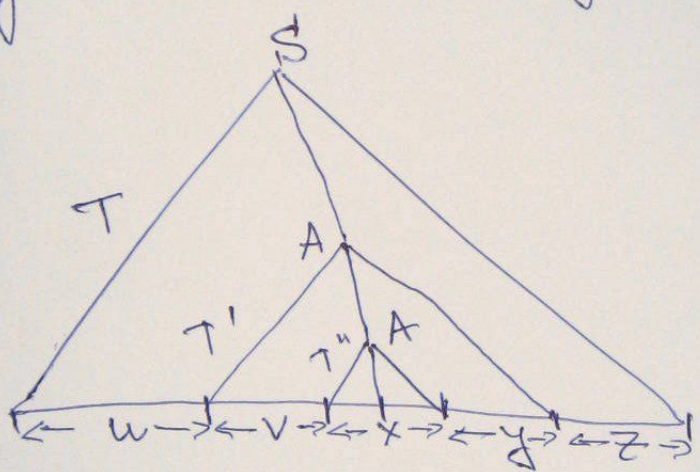
Звезда на Клини: $L(G_1)^*$ се генерира от:

$$G = (V_1 \cup \{S\}, \Sigma_1, R_1 \cup \{S \rightarrow e, S \rightarrow S S_1\}, S)$$

Т Нека $G = (V, \Sigma, R, S)$ е КСГ. Това ва $\forall$ дума $w \in L(G)$ с дължина повече от $\phi(G)^{|V-\Sigma|}$ може да бъде представена във вида $w = uvxyz$, така че или v или y е непразно и $uv^nxy^nz \in L(G)$ за $\forall n \geq 0$.

Доказателство:

Нека w е такава дума и T е дървото на извода с корен S и с ~~резултат yield~~, който има най малък брой нива сред всички дървета на извод със същите корен и ~~резултат yield~~. Той като ~~резултат yield~~ът на T е по-дълъг от $\phi(G)^{|V-\Sigma|}$ следва, че в T има път с дължина поне $|V-\Sigma|+1$, т.е. с поне $|V-\Sigma|+2$ върха. Само един от тези върхове може да е означен с терминал $\Rightarrow$ останалите са означени с не терминали. Той като по пътя има повече върхове отколкото терминали, два върха от пътя са означени със същия елемент A от $V-\Sigma$. Да разгледаме пътя в повее детайли:



~~резултат yield~~ на T''
 v - част от ~~резултат yield~~ на T'
 до началото на ~~резултат yield~~ на T''
 x - част от ~~резултат yield~~ на T
 до началото на T'
 z - останалата част от ~~резултат yield~~ на T .

Оттук е ясно, че частта от T' изключвайки T'' може да бъде повторена произволен брой пъти (включително 0), за да се получат други дървета на извода с ~~резултат yield~~ uv^nxu^nz , $n \geq 0$. Остава само да докажем, че $vy \neq \epsilon$. Но ако $vy = \epsilon$, тогава има дърво с корен S и ~~резултат yield~~ w с по-малко нива от T , получено ако от T премахнем частта от T' , която изключва T'' $\Rightarrow$ противоречие.

Т Контекстно-свободните езици не са затворени относно сечение или допълнение. Доказателство: Виж допълнението.

~~Алгоритъм за КСГ~~

3

D] КСГ $G = (V, \Sigma, R, S)$ е в Нормална форма на Томски ако $R \subseteq (V - \Sigma) \times V^2$.

T] За всяка КСГ $G \exists$ КСГ G' в Нормална форма на Томски, такава че $L(G') = L(G) - (\Sigma^* \{e\})$. Още повече, конструирането на G' може да бъде извършено за полиномиално отношение по отношение на G време.

Доказателство:

Нека $G = (V, \Sigma, R, S)$. Съществуват 3 вида правила $A \rightarrow X$, които не отговарят на условията:

1) Дълги правила ($|X| \geq 3$).
Нека $A \rightarrow B_1 \dots B_n \in R$, $B_1, \dots, B_n \in V$, $n \geq 3$. Заместваме с $n-1$ нови правила:

$A \rightarrow B_1 A_1$
 $A_1 \rightarrow B_2 A_2$
 $A_2 \rightarrow B_3 A_3$
 $\dots$
 $A_{n-2} \rightarrow B_{n-1} B_n$

$A_1, \dots, A_{n-2}$ - нови нетерминали. След прилагане за $\forall$ такова правило получаване еквивалентна граматика.

2) ϵ -правила ($A \rightarrow \epsilon$).

За целта определяме множеството:
 $E = \{A \in V - \Sigma : A \Rightarrow^* \epsilon\}$

Чрез следния алгоритъм:

$E := \emptyset$
while there is rule $A \rightarrow \alpha$, $\alpha \in E^*$ и $A \notin E$ do
add A to E

След това изтриваме от G всички ϵ -правила и: За $\forall$ правило от вида $A \rightarrow BC$ или $A \rightarrow CB$, където $B \in E$ и $C \in V$ добавяме в граматиката $A \rightarrow C$. Единствения извод, който новата граматика би могла да изпусне, е $S \Rightarrow^* \epsilon$ - за щастие, формулирането на теоремата позволява това.

3) Все правила ($A \rightarrow a$ или $A \rightarrow B$)
 За $\forall A \in V$ изчисляване множеството $D(A) = \{B \in V : A \Rightarrow^* B\}$:
 $D(A) := \{A\}$
 while there is a rule $B \rightarrow C$ with $B \in D(A)$ and $C \notin D(A)$ do
 add C to $D(A)$.

Ако a е терминал, то $D(a) = a$.

Премахване всички $\forall$ правила и заместване $\forall$ правила
 от вида $A \rightarrow B^*C^*$ с всички възможни правила от типа
 $A \rightarrow B'C'$, $B' \in D(B)$, $C' \in D(C)$. Накрая добавяме правилата
 $S \rightarrow BC$ за $\forall$ правило $A \rightarrow BC$ такова че $A \in D(S) - \{S\}$.
 Този път е възможно да пропуснем правила от вида $S \rightarrow a$, но
 отново това е позволено от избора на S .

Нещо повече, получените алгоритъм е полиномиален по време относно
 $n = |R|$.

Степка 1) отнема $O(n)$ и остава граматика с големина $O(n)$.
 Втората отнема $O(n^2)$, а третата - $O(n^3)$. С това коремата е доказана.
 Предимството на нормалната форма на Томски е в това, че позволява
 използването на полиномиален алгоритъм за определяне дали
 дума може да се генерира от КСТ:

```

for  $i := 1$  to  $n$  do  $N[i, i] := \{x_i\}$ ; All other  $N[i, j]$  are empty
for  $s := 1$  to  $n$  do
  for  $i := 1$  to  $n-s$  do
    for  $k := i$  to  $i+s-1$  do
      if there is a rule  $A \rightarrow BC \in R$  with  $B \in N[i, k]$  &  $C \in N[k+1, i+s]$ 
        add  $A$  to  $N[i, i+s]$ 

```

Accept x if $S \in N(1, n)$.

Допълнение 5:

8

Примери за езици, които не са контекстно-свободни:

$$\rightarrow L = \{a^n b^n c^n : n \geq 0\}$$

Да допуснем, че $L = L(G)$ за някаква КСГ G . Нека $n > \phi(G)^{1/2-1}$.
 Тогава $w = a^n b^n c^n$ е от $L(G)$ и може да се представи във вида $w = uvxyz$ така че v или y е непразно и uv^mxy^mz е от $L(G)$ за всяко $m=0,1,2,\dots$. Има два случая:
 - vy съдържа и трите символа a, b, c , тогава поне едно от v, y трябва да съдържа поне два от тях (Дирихае). Но тогава uv^2xy^2z съдържа две срещанятия в "грешен" ред - или "b" преди "a", или "c" преди "a" или "b".
 - vy съдържа само някои, но не всичките три символа. Но тогава uv^2xy^2z съдържа неравен брой a, b, c .

$$\rightarrow L = \{a^n : n \geq 1\} \text{ е просто число}$$

Нека $p = \phi(G)^{1/2-1}$, където G е КСГ, която генерира L .
 Тогава $w = a^p$ може да бъде представена във вида $w = uvxyz$, $vy \neq \epsilon$. Нека $vy = a^q$ и $uxz = a^r$, $q, r \in \mathbb{N}$, $q > 0$. Тогава, според теоремата, $r + nq$ е просто, $\forall n \geq 0$. Но това е невъзможно, защото ако например $n = 2q + r + 2 \Rightarrow r + 2q^2 + q(r+2q) = r(q+1) + 2q(q+1) = (2q+r)(q+1) \equiv 2 \cdot 2 = 4$ е просто = 4.

$$\rightarrow L = \{w \in \{a, b, c\}^* : w \text{ има равен брой } a, b, c\}$$

Да допуснем, че L е КС $\Rightarrow$ по Тн за сечение на КС и регулярен език, L ще бъде и сечението на L с регулярен език $a^*b^*c^*$. Но ние току-що доказавме, че този език $(\{a^n b^n c^n : n \geq 0\})$ не е КС.

6 - допълнение:

Контекстно-свободните езици не са затворени спрямо сечение или допълнение.

Доказателство:

Едно е, че $\{a^m b^n c^m : m, n \geq 0\}$ и $\{a^m b^n c^n : m, n \geq 0\}$ са КС. Сечението им е езикът $\{a^n b^n c^n : n \geq 0\}$, който показавме, че не е КС. И тъй като:

$$L_1 \cap L_2 = \overline{L_1 \cup L_2} \text{ и понеже КСЕ са затворени спрямо } \cup,$$

ако КСЕ бяха затворени спрямо отрицание, щяха да бъдат затворени и спрямо сечение (а те не са).
 $\nexists$ КСЕ не са затворени спрямо $\cap$ и $-$.