

1. část: Najděte jednoznačnou grammatiku generující $L = \{a^i b^i \mid i \in \mathbb{N}\}$

$S \rightarrow AS \mid BS \mid A \mid B$

$A \rightarrow aXb$

$B \rightarrow bYa$

$X \rightarrow AX \mid \lambda$

$Y \rightarrow BY \mid \lambda$

Ukomentář:

Pomocí pravidla $S \rightarrow AS \mid BS \mid A \mid B$ můžeme generovat různé permutace A ček a B ček.

Následný přepis A ček a B ček mi zaručuje „párovost“ a, b a zároveň pomocí X, Y dovoluje rozšíření nepřímě o další neterminály / párové terminály.

Důkaz jednoznačnosti:

Dokážu pro první polovinu pravidel $S \rightarrow AS \mid A$, $A \rightarrow aXb$ a $X \rightarrow AX \mid \lambda$.
Pro druhou polovinu je důkaz stejný.

Uvážme, že terminály a, b představují párovky $a: (, b:)$.

Pak je jednoduše vidět, že pravidlo $\rightarrow$ společně s $\rightarrow$ generuje korektní uzávěrování, protože $\rightarrow$ generuje „ (X) “, zatímco $\rightarrow$ generuje „ $(X)(X)\dots$ “, tedy $\rightarrow$ „konečí“, zatímco $\rightarrow$ slouží párovky vedle sebe.

Platí, že pro lib. korektně uzávěrované slovo „ w “ neexistuje žádný vlastní prefix řetězce „ (w) “, který by byl korektně uzávěrovaný. Uvážme pro spor, že je. Pak takový řetězec začneme „loupat“ o uzavřené párovky. Pokud by byl prefix správně uzávěrovaný, měl bych řetězec dlouhý do prázdného slova, avšak ve slově mi zbyde právě ty druhé strany párovky, které jsem vypustil jako suffix. Tedy existuje vždy unikátní korektně uzávěrované slovo.

Nyní ukážu, že i pravidlo $S \rightarrow AS \mid A$ odvozuje unikátně. Pravidlo $S \rightarrow A$ také lze aplikovat pouze na korektně uzávěrované slovo, aby gramatika dávala smysl.

Tudíž na slovo ve tvaru „ (X) “. O takovém slově však již víme, že nemá žádný prefix,

ktej by byl lechtne vzorovan. Tedy speciální by takové slovo nešlo zpracovat pravidlem $S \rightarrow AS$, protože by to bylo ve sporn s existencí lechtného prefixu. Tudiž každé slovo musí mít jednoznačné odvození /strom odvození.

To samé platí adekvatně pro $S \rightarrow BS/B$, $B \rightarrow bYa$, $Y \rightarrow BY$.

Tedy každé slovo má unikátní odvození, protože jsme dokázali, že žádná dvě rozdílná slova se nemohou plně shodovat ve svém odvození.



Dokažte, že následující gramatika generuje reg. jazyk.

2) $X_i \rightarrow a \quad \forall i \in \{1, \dots, n\}$

$X_i \rightarrow X_j X_k \quad \forall j \geq i, \forall k > i, \forall i \in \{1, \dots, n-1\}$

Pro zadanou obecnou gramatiku vytvoříme RegEx, který nám automaticky zaručuje regulárnost jazyka. Důkazem tedy postupně bude korektnost jednotlivých kroků přepisu pro vytvoření RegExu.

Důležitým pozorováním je, že pokud je n končné, je i počet pravidel pro dané $X_i \rightarrow X_j X_k$ $j \geq i, k > i$ končný.

Dále můžeme uvažovat, že $S \rightarrow X_i$ pro lib. $i \in \{1, \dots, n\}$. $\rightarrow$ je také končné množ.

Pak gramatické pravidlo pro X_i lze přepsat do RegExu následovně:

$$X_i \sim \left(\underbrace{\left(X_i (X_{i+1})^+ + X_i (X_{i+2})^+ + \dots + X_i (X_n)^+ \right)}_{n \text{ termínů}} + \underbrace{\left(X_{i+1} X_{i+1} + X_{i+1} X_{i+2} + \dots + X_{i+2} X_{i+1} + \dots + X_n X_n \right)}_{n \text{ termínů}} + (a + b + \dots) \right)$$

Ude seže X_i reprezentuje pravidlo $X_i \rightarrow X_j X_k$ $j \geq i, k > i$. lze uvažovat, že použití tohoto přepisu v RegExu (díky jeho končnosti) implikuje, že bude slovo generováno pouze tímto pravidlem a tedy rozhodně se musí dříve či později, daný nonterminal X_i přepsat na lib. terminal a pomocí pravidla $X_i \rightarrow a$, protože jinak je jen použito „překládání“ pravidla $X_i \rightarrow X_i X_j$, které X_i vymění.

Tedy lze první část přepsat na:

$$\underline{\left((a+b+\dots)(X_{i+1})^+ + (a+b+\dots)(X_{i+2})^+ + \dots + (a+b+\dots)(X_n)^+ \right)}$$

kde $(a+b+\dots)$ je disjunkce n terminálů.

↳ protože 1. pravidlo $X_i \rightarrow a$, kde a je lib. terminál

✏ tuto reprezentuje všechny pravidla $X_i \rightarrow X_j X_k$ $j \geq i, k > i$. Tam žádná opakování nelze zaručit a tudíž vyjde zavazet mnozí žádné zjednodušení.

✏ zde si lze všimnout následujícího: (příklad) $(a(X_j)^+ + a) \sim a(X_j)^*$.

Tedy pokud v $\bullet$ nahradím všechny $^+$ hvězdičkou, můžu celou $\bullet$ část vypustit.

Mám tedy limitu:

$$X_i \sim \left(\left((a+b+\dots)(X_{i+1})^* + (a+b+\dots)(X_{i+2})^* + \dots + (a+b+\dots)(X_n)^* \right) + \left(X_{i+1}X_{i+1} + X_{i+1}X_{i+2} + \dots + X_{i+1}X_n + X_{i+2}X_{i+1} + \dots + X_nX_n \right) \right)$$

Nyní si stačí všimnout, že ve výrazu se vyskytují nejmenší X_{i+1} neterminály a žádné X_i , tedy úplně stejně můžu začít přepisovat všechny X_{i+1} podle stejného algoritmu.

Jelikož je ale neterminál jen konečně mnoho (a protože nemůžeme uvést „overflow“ neterminály) jako např.: X_{n+1} , které z výrazu jednoduše vypustíme), tak po konečné množině kroků v konečné množině iterací získáme RegEx, který bude obsahovat jenom terminály, protože speciálně pro X_n existuje pouze $X_n \rightarrow a$.

Kždý RegEx generuje reg. jazyk a zároveň RegEx generuje jazyk daný gramatikou, protože z ní přímo vychází a její pravidla jen „formalizuje“ do výrazu.

Celkový RegEx tedy bude gigantický, ale protože každá iterace má konečnou množinu kroků a konečnou dlouhý RegExem, i finální výraz bude konečný.

RegEx vytvořený iterativním přepisem pro lib. X_i je tedy důkazem regulárnosti gramatiky.