

Jazyky pro empirii a teorii

KAREL ČULÍK

Je popsán postup vědeckého poznání vedoucího od poznávané skutečnosti k jazyku; přitom je využito základních pojmů predikátové logiky a teorie modelů.

1. URČENOST JAZYKA

Jazyky vznikly jako sdělovací prostředky ve společnosti. Pomocí jazyka někdo sděluje někomu něco, takže povaha a vlastnosti jazyka mohou být ovlivněny jednak tím kdo a komu sděluje, jednak tím co sděluje. V každém případě se sdělují jednotlivé jazykové výrazy, které však mluvčího i posluchače k něčemu odkazují, tj. ke svým významům. Posluchač výrazu mluvčího rozumí správně, když jej odkazuje k témuž významu jako mluvčího. Tedy dříve než se může výrazem něco sdělovat musí mít tento výraz svůj význam, tj. musí k něčemu odkazovat nebo musí něco vyjadřovat. Je tedy tento dvoučlenný vztah odkazování výrazů k jejich významům nejdůležitější pro určení jazyka. Proto jazyk L pokládáme za určen, když je určena množina jeho výrazů E , množina jeho významů M a jeho odkazovací relace $q \subset E \times M$, která každému výrazu $e \in E$ přiřazuje jeden nebo více významů $m \in M$, tj. platí $(e, m) \in q$.

Při tom je třeba doplnit, že odkazovací relace je uchována vždy v nějakém *poznávacím subjektu* S , který je nositelem jazyka a ovládá jej. Z tohoto hlediska existují vždy jen individuální jazyky $L^{(s)} = \langle E, M, q \rangle$, protože se zánikem nositele S zaniká i odkazovací relace q a tím i celý jazyk, i když zřejmě jak E tak M mohou přetrvávat dále. Ale zřejmě samotná množina výrazů E bez odkazovací relace q není ještě jazykem.

2. VÝZNAMY A STRUKTURA SKUTEČNOSTI

Při vzniku jazyka vždy předchází to, co jím chceme vyjadřovat. Ale zcela obecně *jazykem vyjadřujeme skutečnost*, tj. jednotlivými výrazy odkazujeme k jednotlivým

částem skutečnosti. To znamená, že významy jsou jednotlivé části skutečnosti. Potom je samozřejmé, že struktura skutečnosti bude patrně odpovídat i struktura jazykových výrazů nebo – podrobněji řečeno – stupni poznání struktury skutečnosti bude odpovídat struktura výrazů. Je dosti obtížné si představit jiný stupeň poznání skutečnosti než je ten, na který jsme zvyklí.

Současný stav poznání skutečnosti lze charakterisovat *rozlišovací schopností* (a patrně značně neomezenou pamětovou schopností), která dovoluje rozlišovat jednotlivé *předměty*, jednotlivé *vztahy mezi předměty*, jednotlivé *situace*, kdy předměty v nějakém vztahu jsou a kdy ne, dále dovoluje *počítat* předměty a *měřit* některé vztahy, určovat *místo* a *čas* a příp. další speciálnější podmínky.

Je obtížné vymezit co je předmět, ať okamžikový nebo historický, když jej nechceme určovat pomocí jeho vlastností, což ne nezbytné, chceme-li zachycovat změnu.

Ještě obtížnější je vymezit co jsou vztahy mezi předměty příp. ve zvláštním případě, co jsou vlastnosti předmětů. Zde budeme vždy mlčky předpokládat, že každému vztahu odpovídá jistý *rozhodovací* (příp. *měřicí*) *postup*, kterým lze jednoznačně rozhodnout o každých předmětech zda v uvažovaném vztahu jsou nebo nikoli.

Je tedy současný stav poznání skutečnosti charakterizován empirickými kategoriemi předmětu, vlastnosti a vztahu, situace a dalšími kategoriemi místa, času, počtu a míry.

3. ÚPLNÉ POZITIVNÍ EMPIRICKÉ POZNÁNÍ A ROZLIŠOVACÍ MODEL SKUTEČNOSTI

Při poznávání skutečnosti se v každém případě omezujeme jen na nějakou její část a kromě toho na nás (na poznávacím subjektu) záleží, kterých předmětů a kterých vlastností a vztahů si budeme všimát. Vlastní empirické poznávání záleží v aplikaci jednotlivých rozhodovacích a měřicích postupů (odpovídajících vlastnostem a vztahům), na zvolené předměty a v zjišťování situací. Výsledkem tohoto poznání je zjištění, jestli postup vedl k očekávanému (obecně řekneme *pozitivnímu*) *výsledku* nebo ne (tedy vedl k *negativnímu výsledku*). To je v případě rozhodovacích postupů, tj. u *kvalitativního poznání*, zatímco u měřicích postupů tj. u *kvantitativního poznání*, je výsledkem číslo.

Takto se získávají empirické primitivní poznatky, což není nic jiného než „*protokolární věty*“ L. Wittgensteina (zde se pro jednoduchost neuvažují údaje o místě a čase). Souhrn všech empirických primitivních poznatků týkajících se zvolených předmětů, vlastností a vztahů se nazývá *úplným empirickým poznáním* těchto předmětů, vlastností a vztahů. Obvykle se zaznamenávají jen pozitivní primitivní poznatky a pak jejich souhrn se nazývá *úplným pozitivním empirickým poznáním* uvedených předmětů, vlastností a vztahů (dohoda je v tom, že když zde chybí primitivní poznatek týkající se jistého předmětů a jisté vlastnosti, znamená to, že při poznávání byl získán v tomto případě negativní poznatek). Každý empirický výzkum ve kterékoliv vědě je poznáváním v uvedeném slova smyslu.

Výsledky uvedeného empirického poznání lze vyjádřit jazykovými výrazy užívanými v predikátové logice takto: především každému z uvažovaných předmětů dáme jméno, tj. zvolíme symbol, který k němu odkazuje tak, že různé symboly odkazují k různým předmětům; nechť např. $\mathcal{I} = \{a, b, c, d\}$ je množinou všech použitých symbolů. Dále také každé z uvažovaných vlastností a vztahů dáme jméno, tj. symbol, který k němu odkazuje a zase tak, že různé symboly odkazují k různým vztahům; nechť např. $\mathcal{P} = \mathcal{P}_1 \cup \mathcal{P}_2 \cup \dots$, kde $\mathcal{P}_i$ je množinou všech symbolů odkazujících k i -místným vztahům pro $i = 1, 2, \dots$ např. $\mathcal{P}_1 = \{B, W\}$, $\mathcal{P}_2 = \{C\}$ a $\mathcal{P}_i = \emptyset$ pro $i > 2$. Navíc předpokládejme, že symboly z $\mathcal{I}$ a $\mathcal{P}$ dovedeme rozlišit. Konečně situaci, že předmět a má vlastnost B vyjádříme složeným symbolem $B(a)$, tj. $B(a)$ vyjadřuje *pozitivní primitivní poznatek*, že při aplikaci rozhodovacího postupu, který odpovídá vlastnosti B , na předmět a byl výsledek pozitivní, a podobně dostaneme složené výrazy $B(b)$, $W(a)$, $C(a, b)$ atd. Označíme-li souhrn všech pozitivních primitivních poznatků týkajících se předmětů označených symboly z $\mathcal{I}$ a vztahů označených symboly z $\mathcal{P}$ jako $\mathcal{E}(\mathcal{I}, \mathcal{P})$, může být např. $\mathcal{E}(\mathcal{I}, \mathcal{P}) = \{B(a), B(c), W(b), W(d), C(a, b), C(c, b), C(c, d), C(d, b)\}$. Chceme-li vyjádřit že b nemá vlastnost B , musíme zavést nový symbol, např. symbol $\sim$, takže $\sim B(b)$ vyjadřuje *negativní primitivní poznatek*. Podobně je tomu s jinými předměty a vlastnostmi a vztahy.

Byl zde tedy určen jazyk, jímž jsme vyjádřili výsledky empirického poznání, totiž množinou jeho výrazů je množina $\mathcal{I} \cup \mathcal{P} \cup \mathcal{E}(\mathcal{I}, \mathcal{P})$, množinou jeho významů jsou všechny předměty, vlastnosti a vztahy, které jsme uvažovali a všechny situace, v nichž tyto předměty, vztahy a vlastnosti jsou, a relace odkazování byla určena velmi podrobně (i když vlastní akt přiřazení jmen, protože žádné konkrétní předměty, vlastnosti a vztahy nebyly zvoleny, nemohl být předveden).

V logice se obvykle symboly z $\mathcal{I}$ nazývají *individuovými konstantami*, symboly z $\mathcal{P}$ se nazývají *predikátovými konstantami* a složené symboly z $\mathcal{E}(\mathcal{I}, \mathcal{P})$ se nazývají *primitivní výroky*.

Úplně pozitivní empirické poznání je tedy vyjádřeno souhrnem pozitivních primitivních výroků v $\mathcal{E}(\mathcal{I}, \mathcal{P})$.

Odkazovací relace mezi individuovými konstantami a příslušnými předměty se obvykle nazývá *denotací* a místo o významu se v tomto případě hovoří o *denotátu* nebo o *extenzi*. Aby nedošlo k nedorozumění budeme v případě denotace hovořit o *extenzionálním významu*. Odkazovací relace mezi predikátovými konstantami a příslušnými vlastnostmi či vztahy se někdy nazývá *konotací* a místo o významu se v tomto případě hovoří o *intenzi*. Proto budeme v případě konotace hovořit o *intenzionálním významu*. To je nutné proto, abychom od intenzionálního významu nějaké predikátové konstanty, např. dvoumístné predikátové konstanty C , odlišili její *extenzi* či *extenzionální význam vzhledem k dané množině* předmětů $\mathcal{I}$, což je množina všech uspořádaných dvojic předmětů označených symboly x, y , což jsou individuové konstanty z $\mathcal{I}$ takové, že $C(x, y) \in \mathcal{E}(\mathcal{I}, \mathcal{P})$.

Obdobně je určena extenze jiných predikátových konstant.

Množina individuových konstant spolu s extenzemi všech predikátových konstant zavedených pro vyjádření dané části skutečnosti (se zvolenými předměty a vztahy) se nazývá *rozlišovacím modelem* dané skutečnosti. V našem případě je rozlišovací model určen takto: $\langle \mathcal{I} = \{a, b, c, d\}; B = \{a, c\}, W = \{b, d\}, C = \{(a, b), (c, b), (c, d), (d, b)\} \rangle$. Je zřejmé, že rozlišovací model je ekvivalentní úplnému pozitivnímu empirickému poznání téže části skutečnosti, neboť od jednoho lze přejít k druhému zcela formálně. Rozlišovací model je zřejmě množinovým modelem, v němž prvky jsou přímo základní symboly užité jako individuové konstanty, tj. za abstraktní množiny volíme jako zvláštní možný případ přímo množiny symbolů.

4. TEORETICKÝ PŘÍSTUP. LOGICKÉ MOŽNOSTI

Avšak již před aplikací rozhodovacích postupů na nějaké předměty víme, že očekáváme pouze dvě *logické možnosti*. Je-li výsledek pozitivní, pak o výroku, který příslušnou situaci vyjadřuje, říkáme, že je *pravdivý* neboli, že má *pravdivostní hodnotu* 1, a je-li negativní, hovoříme o *nepravdivém* výroku, který má *pravdivostní hodnotu* 0. Tyto dvě logické možnosti a pak i dvě pravdivostní hodnoty jsou dány naším způsobem kladení otázek poznávané skutečnosti (totiž, že my zjišťujeme zda-li předmět vlastnost má nebo nemá) a naším způsobem vyjadřování se o odpovědích na tyto otázky.

Všude tam, kde uvažujeme možnosti, jsme na půdě teorie. Na rozdíl od empirického přístupu v předešlém odstavci, kde např. $B(a)$ bylo pozitivním primitivním výrokiem a $\sim B(b)$ negativním bez ohledu na jakékoliv další možnosti a jen s ohledem na poznávanou část skutečnosti, při teoretickém přístupu se složený symbol $B(a)$, ještě před aplikací rozhodovacího postupu B na předmět a , považuje za *pravdivostní proměnnou*, která může nabývat obou pravdivostních hodnot 1 nebo 0. Tento nepřijemný fakt, že též složený symbol $B(a)$ z hlediska empirického vyjadřuje pozitivní výrok, tj. je jím určena pravdivostní hodnota 1, ale z hlediska teoretického vyjadřuje pravdivostní proměnnou, tj. není jím určena žádná z obou pravdivostních hodnot, je třeba mít vždy na paměti, aby nedošlo k nedorozuměním a záměnám.

Všechny *pravdivostní funkce* výrokové logiky jsou právě funkce libovolného počtu proměnných $X_1, \dots, X_n$, kde $n \geq 1$, jejichž oborem proměnnosti jsou uvedené dvě logické možnosti, tj. $\{0, 1\}$. Proto X_i , $1 \leq i \leq n$ se rovněž nazývá pravdivostní proměnnou. Oborem hodnot pravdivostních funkcí jsou obě pravdivostní hodnoty 0, 1. Pravdivostní funkce $f(X_1, \dots, X_n)$ může být zadána tabulkou nebo nějakým booleovským výrazem. Obvykle jsou booleovské výrazy definovány rekurentně z jednodušších, když se na základní pravdivostní funkce volí disjunkce „ $\vee$ “, konjunkce „ $\&$ “, negace „ $\sim$ “, implikace „ $\rightarrow$ “ a ekvivalence „ $\leftrightarrow$ “, kterým v přirozeném jazyce odpovídají známé spojky a rčení „buď a nebo“, „i“, „není pravda, že“, „jestliže, pak“ a „tehdy a jen tehdy“. Je dobře známé co jsou to tautologie a kontradikce i jak se dají booleovské výrazy srovnávat.

Tedy např. $X_1 \vee X_2$ je booleovský výraz, ale také $B(a) \vee B(b)$ nebo $W(a) \vee$

$\vee C(b, c)$ jsou booleovské výrazy, neboť X_1, X_2 i $B(a), B(b), W(a), C(b, c)$ jsou pravdivostní proměnné. Při tom je zcela druhotné, že jak „B“ tak „a“ atd. mají své pevně určené významy, a že výrazu $B(a)$ proto rozumíme (aniž víme, zdali je pravdivý nebo nepravdivý). Tato okolnost je naopak zavádějící a důvodem rozpaků, které pocítujeme při zcela konkrétních příkladech, jako např. „Buď Praha je hlavní město Československa nebo Berlín leží na Seině“, kde totiž obě věty nejsou pravdivostními proměnnými, nýbrž konstantami, protože jich užíváme z hlediska empirického. Pak ovšem uvedená věta nevyjadřuje pravdivostní funkci tj. disjunkci jako třeba $X_1 \vee X_2$, nýbrž místo toho jenom $1 \vee 0$, což disjunkci ovšem předpokládá.

Posloupnosti z 0 a 1 délky $k \geq 1$ nazvěme *k-násobnými logickými možnostmi*. Např. pro $k = 2$ jsou celkem 4 2-násobné logické možnosti a to $(0, 0), (0, 1), (1, 0), (1, 1)$. Obdobně, jako se dá každá pravdivostní funkce vyjádřit svou úplnou normální disjunktivní formou, lze každou pravdivostní funkci k proměnných vyjádřit pomocí množiny všech k -násobných logických možností, pro něž nabývá pravdivostní hodnoty 1.

Tedy např. $X_1 \vee X_2$ je vyjádřena takto: $[X_1, X_2; (1, 1), (1, 0), (0, 1)]$ apod. Timto způsobem lze, podobně jako u tabulky, vyhnout se užívání logických spojek.

5. SLOŽENÉ PREDIKÁTY

Z hlediska empirického poznání – pokud jsme se omezili na množinu předmětů $\mathcal{I}$ a množinu vlastností a vztahů $\mathcal{P}$ – nelze již k $\mathcal{E}(\mathcal{I}, \mathcal{P})$ nic více připojit. To je také důvod proč se $\mathcal{E}(\mathcal{I}, \mathcal{P})$ nazývá úplným poznáním. Jakékoliv další empirické poznatky lze získat jenom buď přibráním dalších předmětů nebo dalších vlastností a vztahů, ale těm bude zase odpovídat další úplné pozitivní empirické poznání atd. v závislosti na tom, co se poznávací subjekt rozhodne zkoumat. Jakým směrem může pokračovat poznání za $\mathcal{E}(\mathcal{I}, \mathcal{P})$ pro zvolenou část skutečnosti? Co lze dále poznávat na zvolené části skutečnosti, když jsme empirické poznání získáním $\mathcal{E}(\mathcal{I}, \mathcal{P})$ již ukončili. Neformálně tato otázka zní: co se může dále dělat, tj. jak zpracovávat získané empirické údaje?

Přestože $\mathcal{I}$ a $\mathcal{P}$ jsou pevně zvoleny a $\mathcal{E}(\mathcal{I}, \mathcal{P})$ je určeno, lze v poznávání zvolené části skutečnosti pokračovat. Zavedeme nové vlastnosti a vztahy tím, že pomocí pravdivostních funkcí sestavujeme z rozhodovacích postupů pro vlastnosti a vztahy z $\mathcal{P}$ rozhodovací složené postupy pro složené vlastnosti a vztahy. Které složené vlastnosti a vztahy budeme zkoumat, to záleží na nás, tj. na poznávacím subjektu.

Individuovou proměnnou nad $\mathcal{I}$ nazveme každý symbol x , o němž přijmeme dohodu, že x může označovat kterýkoliv z předmětů označených individuovými konstantami z $\mathcal{I}$. Tedy množina předmětů označovaná symbolem $\mathcal{I}$ je *obor proměnnosti proměnné x* . Opět se samozřejmě žádá, aby se daly odlišit proměnné od konstant. Jestliže v primitivním výroku $B(a)$ nebo $C(b, c)$ nahradíme individuovou konstantu (jednu nebo více) individuovou proměnnou, vznikne *schéma* $B(x)$ nebo $B(y)$

nebo $C(x, y)$ nebo $C(x, x)$ atd. Zřejmě počet různých schémat nezáleží jen na počtu základních predikátů, ale i na počtu různých individuových proměnných, které můžeme používat.

Vezmeme-li nyní libovolný booleovský výraz v pravdivostních proměnných $X_1, \dots, X_n$ a dosadíme-li za tyto proměnné nějaká schémata, pak získané *složené schéma* vyjadřuje tolikamístný složený vztah, kolik různých individuových proměnných se ve složeném schématu vyskytuje. Např. z $X_1 \vee X_2$ lze dostat složené schéma $B(x) \vee W(x)$ (ale také třeba $B(x) \vee W(y)$ nebo $B(x) \vee B(y)$ apod.) vyjadřující zřejmě zase vlastnost složenou uvedeným způsobem z obou základních vlastností a její rozhodovací postup je zřejmě tento: vezmou se oba rozhodovací postupy pro B i W a složený postup dá pozitivní výsledek právě tehdy, když alespoň jeden z postupů pro B a W dá pozitivní výsledek. Snadno je vidět jak se skládají rozhodovací postupy ve složitějších případech.

Důležité pro složené rozhodovací postupy je, že – pokud $\mathcal{E}(\mathcal{I}, \mathcal{P})$ je známo – není vůbec potřeba jednotlivé základní rozhodovací postupy skutečně na jednotlivé předměty aplikovat, ale že stačí se dívat na $\mathcal{E}(\mathcal{I}, \mathcal{P})$ a podle toho jen vyhodnocovat příslušnou booleovskou funkci. Tedy např. pro uvedený složený predikát $B(x) \vee W(x)$ a pro předmět $a \in \mathcal{I}$ zjistíme, že $B(a) \in \mathcal{E}(\mathcal{I}, \mathcal{P})$, ale $W(a) \notin \mathcal{E}(\mathcal{I}, \mathcal{P})$, tj. bylo $\sim W(a)$, což z teoretického hlediska přepsáno znamená, že $B(a)$ má hodnotu 1, zatímco $W(a)$ má hodnotu 0, takže $B(a) \vee W(a) = 1 \vee 0 = 1$. Je tedy $B(a) \vee W(a)$ pozitivním složeným výrokem, čili je pravdivá, ale nikterak se to na ní nepozná, dokud neuvedeme, že $B(a)$ je pravdivá a $W(a)$ není pravdivá. To je důsledek teoretického přístupu, kde $B(a)$ i $W(a)$ jsou pravdivostní proměnné a nikoli pravdivé či nepravdivé výroky.

Dalším důležitým faktem je, že pozitivní složené výroky, např. $B(a) \vee W(a)$ a to včetně doplňku, že $B(a) = 1$, $W(a) = 0$, nevyjadřující žádnou situaci v poznávané části skutečnosti se zvoleným $\mathcal{I}$ a $\mathcal{P}$. Tato okolnost vynikne, když se vyhneme užívání logických spojek. Pak složený predikát $B(x) \vee W(x)$ je vyjádřen jako $[B(x), W(x); (1, 1), (0, 1), (1, 0)]$ a tedy pravdivost složeného výroku $B(a) \vee W(a)$, když $B(a) = 1$ a $W(a) = 0$, znamená pouze to, že skutečně nastala jedna z 2-násobných logických možností, které byly předepsány.

6. OBECNÉ POZNATKY. DEDUKCE

Vlastní smysl zavádění složených vlastností a vztahů se objeví teprve když uplatníme další důležitou charakteristiku našeho poznávání skutečnosti, totiž snahu po *obecných poznatcích*, tj. o celých skupinách primitivních (nebo složených) výroků, které se všechny týkají téhož predikátu (nebo složeného predikátu) a všech uvažovaných předmětů.

Jestliže $R(x)$ je schéma odpovídající základnímu nebo složenému predikátu vzhledem k $\mathcal{E}(\mathcal{I}, \mathcal{P})$, tj. buď např. $R(x) = {}_{at}B(x)$ (nebo např. $R(x) = {}_{at}B(x) \vee W(x)$), pak chceme zjišťovat, zda-li $R(x)$ je pravdivé pro každé $x \in \mathcal{I}$, tj. chceme zjišťovat

pravdivost výroků $R(a)$, $R(b)$, $R(c)$, $R(d)$, když v našem případě $\mathcal{J} = \{a, b, c, d\}$. Rovněž toto zjištění lze získat bez dalšího zkoumání dané skutečnosti, neboť opět stačí jen si všimnout $\mathcal{E}(\mathcal{J}, \mathcal{P})$, tedy zase jistým teoretickým poznávacím postupem. Jestliže všechny uvedené jednotlivé výroky jsou pravdivé, pak místo nich jako zkratku píšeme *obecný výrok* $\forall R(x)$, a jestliže alespoň jeden z nich je pravdivý, píšeme *existenční výrok* $\exists R(x)$, který ovšem není (zatím zde) zkratkou jako u obecného výroku.

Uvedené kvantory čteme obvyklým způsobem.

Z empirického hlediska $\forall R(x)$ je pravdivý a $\sim \forall R(x)$ je nepravdivý obecný výrok a každý z nich může vyjadřovat jistou složenou situaci popsanou výše. Naproti tomu z teoretického hlediska je $\forall R(x) =_{df} R(a) \& R(b) \& R(c) \& R(d)$, což je jistá pravdivostní funkce, takže celý výraz $\forall R(x)$ je pouze pravdivostní proměnnou.

Např. $\forall B(x)$ je v našem případě nepravdivé, ale $\forall (B(x) \vee W(x))$ je pravdivé, jak se snadno přesvědčíme nahlédnutím do $\mathcal{E}(\mathcal{J}, \mathcal{P})$.

V obecném případě vztahu (nebo složeného vztahu) se z odpovídajícího schématu $R(x_1, \dots, x_n)$ dá vytvořit celá řada obecných či existenčních výroků, např. $\forall \dots \forall R(x_1, x_2, \dots, x_n)$ nebo $\exists \dots \exists R(x_1, x_2, \dots, x_n)$ atd. až $\exists \dots \exists \dots \exists R(x_1, x_2, \dots, x_n)$. Samozřejmě, že se o pravdivosti všech takových obecných výroků — vesměs se stále týkajících zvolené části skutečnosti $(\mathcal{J}, \mathcal{P})$ — můžeme přesvědčit nahlédnutím do $\mathcal{E}(\mathcal{J}, \mathcal{P})$ a vyhodnocením potřebných booleanových funkcí, i když tato možnost bude v mnoha případech jenom teoretická, protože by zabrala příliš mnoho času i na rychlých počítačích strojích.

Množina všech obecných (tj. utvořených pomocí obecných nebo existenčních kvantorů) výroků, které jsou pravdivé vzhledem k $\mathcal{E}(\mathcal{J}, \mathcal{P})$ se může nazvat úplným *pozitivním teoretickým poznáním* části skutečnosti zadané pomocí $(\mathcal{J}, \mathcal{P})$ a budeme ji označovat $\mathcal{T}(\mathcal{J}, \mathcal{P})$.

Avšak když už jsme některé obecné výroky z $\mathcal{T}(\mathcal{J}, \mathcal{P})$ získali, tj. uvedeným teoretickým postupem vzhledem k $\mathcal{E}(\mathcal{J}, \mathcal{P})$, pak se může stát, že další není třeba získávat tímto zdoluhavým postupem, ale lze je získat teoretickým postupem zvaným *deduktivní dokazování* podle známých pravidel. Možnost takového dokazování je patrně jeden z důvodů proč se vůbec snažíme získávat obecné poznatky.

Např. v našem příkladu se snadno přesvědčíme, že do $\mathcal{T}(\mathcal{J}, \mathcal{P})$ patří následující obecný výrok

$$(aS) \quad \forall_{x \in \mathcal{J}} \forall_{y \in \mathcal{J}} (C(x, y) \rightarrow (\sim C(x, x))) ,$$

což je známá podmínka asymetričnosti binární relace, která v přepisu bez užití logických spojek, ale s předepsáním 2-násobných logických možností vypadá takto:

$$(aS) \quad \langle \forall_{x \in \mathcal{J}} \forall_{y \in \mathcal{J}} [(C(x, y), C(y, x)); \{(1, 0), (0, 1), (0, 0)\}] \rangle .$$

Jestliže však už víme, že podmínka (aS) je v $\mathcal{E}(\mathcal{I}, \mathcal{P})$ splněna, pak se z ní dá dokázat např. podmínka (iR)

$$(iR) \quad \forall_{x \in \mathcal{I}} (\sim C(x, x)) \text{ neboli } \langle \forall; [C(x, x)]; \{0\} \rangle.$$

Při deduktivním dokazování se z obecných výroků odvozují jiné obecné výroky. Naproti tomu teoretický postup, kterým jsme — odvoláním se na $\mathcal{E}(\mathcal{I}, \mathcal{P})$ — získávali pravdivé obecné výroky, protože se při něm jde od primitivních výroků k výrokům obecným, se nazývá *induktivním dokazováním*.

7. TEORIE A JEJÍ MODELY

Všude v předešlých odstavcích $\mathcal{E}(\mathcal{I}, \mathcal{P})$ a $\mathcal{T}(\mathcal{I}, \mathcal{P})$ samozřejmě záviselo na $\mathcal{I}$ a $\mathcal{P}$, takže při zkoumání jiné části skutečnosti, tj. při volbě jiných předmětů, vlastností a vztahů, bychom dostali jiné úplné pozitivní empirické a teoretické poznání. Je ovšem ihned zřejmé, že jsme mohli i při zkoumání jiné části skutečnosti zase množinu předmětů označit symbolem $\mathcal{I}$ a množinu vlastností a vztahů symbolem $\mathcal{P}$ a dokonce i dřívějších individuových konstant a proměnných lze znovu užít (ovšem budou odkazovat k jiným předmětům a vztahům), případně podle potřeby lze je doplnit o další symboly, takže po formální stránce, tj. co do tvarů jednotlivých výrazů, se vůbec nepozná, o kterou část skutečnosti jde (neboť to je určeno teprve významy individuových konstant a predikátových konstant). Tato skutečnost je důvodem proč se buduje jenom jednou celá predikátová logika za předpokladu, že $\mathcal{I}$ i $\mathcal{P}$ obsahují dostatečný počet prvků, tj. budoucích individuových a predikátových konstant. Obvykle se navíc individuové konstanty vylučují, takže $\mathcal{I}$ zůstává jen jako symbol, který *bude označovat* obor proměnnosti individuových proměnných. Při tom samotné $\mathcal{I}$ je proměnná, která může označovat kteroukoliv množinu předmětů, a i symboly z $\mathcal{P}$ jsou predikátové proměnné, které mohou označovat kterékoliv vztahy s příslušným počtem míst. Pak ovšem všechny výrazy, které formálně vypadají jako výroky, nejsou žádnými výroky, tedy nejsou ani pravdivé ani nepravdivé, ale jsou to jen schémata výroků. Tato schémata se stanou výroky když *interpretujeme* $\mathcal{I}$ a $\mathcal{P}$, tj. určíme význam pro $\mathcal{I}$ a pro všechny symboly z $\mathcal{P}$, čili když určíme o které předměty, vlastnosti a vztahy jde. Tedy když vše vztáhneme k nějaké skutečnosti.

Dokazování je však teoretický postup týkající se pouze tvaru uvažovaných výroků či schémat (neboť co do tvaru mezi nimi není rozdíl). Proto se v tomto oboru schémat dá zavést *teorie* jako množina (schémat) obecných výroků uzavřená vůči deduktivnímu dokazování, tj. obsahující všechny důsledky (tj. správná schémata důsledků). Nejdůležitější je případ *axiomatické teorie*, což je teorie, jejíž všechny prvky plynou z daných (schémat) obecných výroků nazývaných *axiomy*. Jsou-li $A_1, \dots, A_n$ dané axiomy, pak nechť $T(A_1, \dots, A_n)$ značí teorii s axiomy $A_1, \dots, A_n$.

Množina předmětů $\mathcal{I}$ a množina vlastností a vztahů $\mathcal{P}$, čili interpretace v hoření slova smyslu, se nazývá *modelem teorie* $T(A_1, \dots, A_n)$, když $A_1, \dots, A_n$ jsou pravdivé

výroky vůči $\mathcal{S}(\mathcal{I}, \mathcal{P})$, čili když $A_i \in \mathcal{T}(\mathcal{I}, \mathcal{P})$, pro $i = 1, 2, \dots, n$. Potom zřejmě také platí $T(a_1, \dots, a_n) \subset \mathcal{T}(\mathcal{I}, \mathcal{P})$. Model se nazývá *extenzionální* nebo *intenzionální* podle toho, jestli interpretace predikátových proměnných z $\mathcal{P}$ je extenzionální nebo intenzionální.

Z hlediska konzistence axiomatických teorií stačí uvažovat jenom extenzionální modely. Ale naopak z hlediska využití axiomatických teorií, např. při odvozování obecných výroků z $\mathcal{T}(\mathcal{I}, \mathcal{P})$ pro zadanou konkrétní část skutečnosti $(\mathcal{I}, \mathcal{P})$, jsou cenné právě modely intenzionální. Možnost takového využití axiomatické teorie se opírá o jednu z hlavních vět, totiž že ta schémata obecných výroků, která se dají dokazovat z axiomů, jsou pravdivými výroky ve všech modelech dané teorie.

Je třeba zde jen upozornit na to, že obvykle se nedá deduktivně dokazovat, že nějaká část skutečnosti $(\mathcal{I}, \mathcal{P})$ je modelem dané teorie, ale naopak, že se to musí dokazovat induktivně a to může být často prakticky neproveditelné. Z toho plyne, že ne vždy musí být nejvhodnější cesta rozvíjení dalších poznatků z $\mathcal{T}(\mathcal{I}, \mathcal{P})$ ta, která vede k budování axiomatických teorií $\mathcal{T}(A_1, \dots, A_n)$, jichž je $(\mathcal{I}, \mathcal{P})$ modelem.

8. JAZYKY V EMPIRICKÝCH VĚDÁCH

Hlavním cílem empirických věd a patrně hlavním cílem poznání vůbec je možnost předvídání. Právě předvídáním v čase se stává vědění a věda vůbec skutečnou mocí, které se dá využít i zneužít. Proto se nelze omezovat na danou množinu předmětů označenou symbolem $\mathcal{I}$, ale je třeba připustit, že je neomezená do budoucnosti. Proto se nezadáva výčetem, ale podmínkou, kterou každý předmět musí splňovat. A je zřejmé, že zde všechny obecné výroky jsou vírami do budoucna a skutečně pravdivými výroky jsou jen pro část předmětů patřících minulosti. Dalším častým omezením je, že vlastně každá vlastnost má svůj obor přípustných předmětů, takže libovolné skládání predikátů nemusí mít vždy smysl. Konečně často se ukazuje, že je třeba vyšetřovat mnohem větší počet predikátových konstant než je to zatím v matematických teoriích obvyklé.

Samozřejmě se zde omezujeme na ty části vědy, které nevyžadují především kvantitativních údajů nebo statistického zpracování, ale jimž stačí zpracování kvalitativní. Tomuto kvalitativnímu pojetí také odpovídá současná logika, která se stále vyhýbá přijetí jednak kvantitativních údajů, jednak soustavného obohacení o číselný parametr, který by vyjadřoval čas a dovolil hovořit o změně. Nicméně axiomatizace fyzikálních teorií ukazují, že i po uvedeném doplnění jazyky budou pořád mít ty základní tvary, které byly popsány výše. Mnohem méně je jasné, jak mají vypadat jazyky v případech, kdy je nezbytný a žádoucí přístup statistický a pravděpodobnostní, jakkoli v matematice samotné právě tyto obory mají stále živý styk s empirickými podněty.

(Došlo dne 14. června 1968.)

Languages for Empirism and Theory

KAREL ČULÍK

The cognizing procedure in scientific research is described which leads from the reality to the language, where the basic notions of predicate logic and model theory are applied.

Doc. Dr Karel Čulík, DrSc., Matematický ústav ČSAV, Žitná 25, Praha 1.