

K problému výběru porovnatelné skupiny podniků pro ocenění a nástin budoucího výzkumu s využitím strojového učení[#]

Veronika Staňková* – Miloš Mařík**

Úvod

V oceňování podniku je porovnání s konkurenčními podniky jednou ze základních využívaných metod. Teorie říká, že firmy musí být porovnatelné především v aspektu výnosnosti, rizika a tempa růstu (Damodaran, 2002). Akademici se ale již neshodují, zda příslušnost ke stejnému odvětví je, nebo není dostatečně spolehlivou proxy proměnnou k uvedeným aspektům pro výběr porovnatelné skupiny (k tomu diskuze Alford 1992, Cheng a McNamara 2000 versus Bhojraj a Lee 2002, Overgaard Knudsen 2017 shrnutá v tomto textu).

Metodologický přístup výběru porovnatelných podniků pro ocenění tržním srovnáním každopádně není sjednocen. Neexistence uznávaných zásad dává příliš prostoru subjektivnímu vlivu oceňovatele, který může vést k nekonzistentním oceněním.

Současná praxe primárně vychází z odvětvového zařazení. Použití odvětvových průměrů může, jak detailněji rozebírám v tomto článku, být dobrý sluha, ale špatný pán. Navíc ani odvětvová klasifikace není jednotná (viz koexistence klasifikačních systémů Standard Industrial Classification (SIC), North American Industry Classification System (NAICS), Global Industry Classification Standard (GICS) apod.), což může ztěžovat správné zařazení podniku. Další komplikace jsou zejména u firem na rozhraní oborů a multioborových podniků. Většina klasifikačních systémů dále není schopna flexibilně reagovat např. na nově vznikající odvětví. Z uvedených důvodů docházím v této práci k závěru, že porovnatelnou skupinu je třeba sestavit na základě podobnosti klíčových podnikových fundamentů.

Cílem tohoto článku je syntéza literatury k tématu výběru porovnatelné skupiny a nástin budoucího výzkumu, který by otestoval potenciál využití strojového učení aplikovaného na výběr porovnatelných podniků pro tržní ocenění. Nicméně ani v případě, že by zamýšlený výzkum potvrdil vhodnost strojového učení, nelze očekávat, že by tento jediný výzkum stačil na to, aby se strojové učení stalo dominantní metodou pro sestavení porovnatelné skupiny v rámci oceňovacího postupu v České republice ještě po mnoho let. Spíše předpokládám, že pozitivní výsledek by znamenal pro teorii začátek intenzivnějšího bádání v této oblasti a pro oceňovací praxi možnost dodatečného úhlu pohledu na výběr porovnatelné skupiny.

Současná teorie a praxe

Výběr porovnatelné skupiny podniků v kontextu ocenění tržním porovnáním

Mezinárodní studie ukazují, že v praxi oceňování podniku (myšleno v širším smyslu než jen znalecké oceňování) dominuje použití metody tržního porovnání (k tomu např. tyto

[#] Článek je zpracován jako jeden z výstupů výzkumného projektu IG104020, který je realizován na Fakultě financí a účetnictví VŠE Praha.

* Ing. Veronika Staňková, doktorand, Katedra financí a oceňování podniku VŠE Praha.

** Prof. Ing. Miloš Mařík, CSc., Katedra financí a oceňování podniku VŠE Praha, ředitel Institutu oceňování majetku VŠE Praha.

výzkumy: Demirakos a kol., 2004; Asquith a kol. 2005, Nel 2010, Pinto a kol. 2015, Stowe 2015 podle: Evergaard Knudsen a kol., 2017, str. 85), a to i v českém prostředí podle průzkumu Vydržel a Soukupová (2012). Jednoznačnou výhodou ocenění tržním porovnáním je jeho přímočarost a rychlost zpracování (Plenborg a Pimental, 2016, str. 1).

Nevýhodou tržního porovnání je, že ve srovnání s výnosovým oceněním, kde je nutné explicitně prognózovat budoucí vývoj hlavních charakteristik podniku, je metoda tržního porovnání jednodušší pouze zdánlivě, neboť hodnota podniku je nakonec určována stejnými parametry (generátory hodnoty), jak se dá dokázat rozvedením jednotlivých násobitelů na dílčí fundamenty s využitím Gordonova modelu (např. Mařík 2011, str. 315-316). V tržním porovnání jsou tyto parametry skryté ve volbě porovnatelné skupiny podniků (dále rovněž PG, z angl. peer group), z níž je odvozen oceňovací násobitel. Metoda tržního porovnání je založena na představě relativní rovnosti ceny substitutů (Evergaard Knudsen a kol. 2017, str. 1) vyplývající ze zákona jedné ceny. Kvalita ocenění je nakonec přímo úměrná kvalitě výběru podniků použitých pro srovnání. Proto Plenborg a Pimental (2016) jmenují výběr porovnatelné skupiny podniků jako první úskalí při implementaci metody tržního porovnání.

V současné literatuře lze identifikovat zhruba tři přístupy (Plenborg a Pimental, 2016, str. 2). Stanovení porovnatelné skupiny může vycházet primárně z (i) odvětvové klasifikace, (ii) podobnosti klíčových podnikových fundamentů (generátorů hodnoty) anebo (iii) využití alternativních metod výběru.

Agregované skupiny podle odvětvové klasifikace

Odvětvová klasifikace se zakládá na přesvědčení o vysoké míře vnitroskupinové podobnosti ekonomických fundamentů v rámci odvětví – neboť lze očekávat, že podniky v jednom odvětví čelí podobnému riziku a růstu zisku (Alford, 1992, str. 95). Alford je v literatuře obecně zmiňován jako základní zdroj tohoto přístupu (např. Plenborg a Pimental, 2016, str. 2 a Janda, 2019, str. 42). Alford (1992) zkoumal přesnost P/E násobitele při odhadu ceny za akcii ve srovnání s aktuální cenou v závislosti na odvětvovém zařazení, rizikovitosti (měřené bilanční sumou (dále též TA, z angl. total assets) a zadlužeností) a očekávaném růstu (měřeno jako rentabilita vlastního kapitálu (dále též RoE) a predikcí růstu dle databáze IBES). Výsledky Alfordova výzkumu ukázaly, že nejvyšší přesnosti se dosáhne buď (i) samotným kritériem odvětvového zařazení, nebo (ii) společným využitím kritérií TA + RoE (Alford, 1992, str. 103), nicméně ani jeden z nich samostatně ani v kombinaci s odvětvím nevede k vyšší přesnosti (Alford, 1992, str. 107), což dle Alforda (1992, str. 99) dokazuje, že odvětví efektivně zachycuje rozdíly v rizikovitosti (měřeno TA) a růstu mezi podniky (měřeno RoE). Podobnou studii provedli Cheng a McNamara (2000), kterou víceméně potvrdili zjištění Alforda. Nejeefektivnější metody výběru porovnatelné skupiny podniků jsou podle autorů kombinace odvětví a RoE a samotné odvětvové zařazení (Cheng a McNamara, 2000, str. 367).

Jiní autoři však poukazují na určitá omezení odvětvové klasifikace – např. neexistence jednoho univerzálního klasifikačního systému. Bhojraj a kol. (2003) provedli výzkum, ve kterém srovnávali čtyři různé systémy odvětvové klasifikace: SIC (Standard Industrial Classification), NAICS (North American Industry Classification System), GICS (Global Industry Classification Standard) a klasifikaci autorů Fama&French. Na základě jejich výzkumu vyšla nejlépe klasifikace GICS, která podle autorů významně lépe vysvětluje cross-sekční variabilitu v oceňovacích násobitelích a dalších finančních ukazatelích (Bhojraj a kol., 2003, str. 770). Dále tito autoři poukázali na to, že GICS se s ostatními schémata shodne v pouhých 56 % (Bhojraj a kol., 2003, str. 747), takže opravdu záleží i na výběru systému klasifikace.

Další výzkumy pak poukazují na problémy související s tím, že vzhledem k dynamickému vývoji ekonomiky dochází na straně firem k přechodům mezi různými odvětvími, někdy i k obtížím analytiků přiřadit správné odvětví (aktivita na pomezí, víc aktivit zároveň apod.). Z pohledu odvětví dochází v současnosti k vytváření nových podnikatelských oblastí, a tak statické odvětvové klasifikace nemusejí dobře odpovídat skutečnosti. Tento problém někteří výzkumníci řeší inovativními přístupy, které detailněji rozvádím v podkapitole Alternativní metody (např. Hoberg a Phillips 2010 a 2016, Chong a Zhu 2012, Fang a kol. 2013, Lee a kol. 2015 a Yang a kol. 2019).

V praxi českých znalců jako zdroj dat převládá dataset odvětvových násobitelů publikovaných prof. Damodaranem – pravděpodobně díky volné dostupnosti na jeho webových stránkách a pravidelné aktualizaci dat. Vzhledem ke zjištění Bhojraj a kol. (2003) o nestejnosti klasifikačních systémů jsem porovnála společnosti publikované v jednotlivých odvětvích dle Damodarana se zařazením podle GICS, který je podle Bhojraj a kol. (2003) nejvhodnější ze čtyř zkoumaných schémat pro oceňování podniku. Dataset publikovaný Damodaranem obsahuje přes 44 tisíc společností rozdělených do 94 odvětví, které lze popsat následovně:

- 15 % odvětví dle Damodarana má ekvivalentní zařazení v klasifikaci GICS (tzv. GICS Industry publikované v databázi Capital IQ (dále též CIQ));
- 66 % odvětví dle Damodarana byla převážně srovnatelných s GICS klasifikací (tj. v průměru 2 % společností v každém odvětví by byla podle GICS zařazena jinak) a
- 19 % odvětví dle Damodarana byla vyhodnocena jako smíšená odvětví (tj. k jednomu odvětví u Damodarana lze přiřadit několik odvětvových kategorií dle GICS).

Dataset publikovaný Damodaranem dále umožňuje rozlišit kromě jednotlivých odvětví, také geografické zařazení¹. Z pohledu českého valuátora je zajímavá otázka, kde hledat Českou republiku (dále též ČR): Damodaran zařazuje ČR do geografické skupiny (rozvinutá) Evropa. Nicméně země jako Slovensko a Slovinsko jsou považovány za Emerging Markets, takže nejsou součástí vzorku pro ocenění společnosti z ČR. Naopak jsou v Damodaranově datasetu pro rozvinutou Evropu zohledněné společnosti např. z Turecka. Pokud bychom oceňovali slovenskou společnost, dle Damodarana bychom měli hledat ve skupině Emerging markets, kde jsou pohromadě nerozvinuté trhy z celého světa, včetně např. země z afrického Sahelu a střední Ameriky. Dále jsem v geografických kategoriích našla i zcela zjevnou chybu, kdy v podskupině Eastern Europe & Russia (obsahuje 17 zemí převážně z bývalého Sovětského svazu) se nachází rovněž Uganda. Považuji tedy za důležité, aby oceňovatel nepřebíral tato data automaticky, ale měl by relevantní skupinu sám aktivně projít a zhodnotit, jestli jednotlivé zahrnuté společnosti dávají smysl pro daný oceňovací případ. Přehled srovnání geografických kategorií u Damodarana a v databázi CIQ ukazuje Tabulka 1.

¹ V Damodaranově datasetu se rozlišuje na 5 základních geografických celků (tzv. Broad Group zahrnující USA; Austrálie, Nový Zéland & Kanada, Evropa, rozvíjející se trhy (dále též Emerging Markets) a Japonsko).

Tabulka 1 Srovnání geografických kategorií u Damodarana a CIQ databázi

Damodaran		
	Název skupiny	Celkem zemí ve skupině
	US	1
	Australia, NZ & Canada	3
	(developed) Europe	30
	Emerging Markets	101
	Japan	1
Celkem:	5 skupin	136 zemí
CIQ		
	Název skupiny	Celkem zemí ve skupině
	United States and Canada	2
	European Developed Markets	27
	European Emerging Markets	22
	Asia / Pacific Developed Markets	6
	Asia / Pacific Emerging Markets	44
	Latin America and Caribbean	45
	Africa/Middle East	67
Celkem	7 skupin (*)	213 zemí

(*) Pozn.: Pro lepší porovnatelnost obou datasetů jsem z databáze CIQ uvažovala nejobecnější členění zahrnující 5 skupin (pouze geografické členění) rozšířené o rozlišení na emerging a developed všude, kde bylo dostupné (CIQ uvádí pouze u regionů Asia/Pacific a European).

Zdroj: vlastní zpracování na základě dat Damodaran a databáze CIQ

Hledání fundamentálních ukazatelů

Vzhledem k nedostatkům běžných odvětvových klasifikačních systémů uvedených výše poukazují někteří výzkumníci, že příslušnost ke stejnému odvětví není zárukou stejných fundamentálních charakteristik (generátorů hodnoty podniku) zaručující podobnost budoucích peněžních toků plynoucích z podniku. Dittmann a Weiner (2005) provedli podobný výzkum jako Alford (1992), ale došli k výrazně jiným závěrům. Na datech z USA a evropského akciového trhu prokázali, že rentabilita aktiv (dále též RoA) anebo kombinace RoA + TA překonává porovnatelné skupiny sestavené pouze na základě příslušnosti ke stejnému odvětví.

Zásadní problém tohoto přístupu je, že při použití hranice 14% nejpodobnějších společností (např. Alford (1992), Dittman a Weiner (2005)) je k tvorbě porovnatelné skupiny podniků při metodě průniku dvou faktorů k dispozici pouze 2% trhu ($14\% \times 14\% = 1,96\%$), při třech faktorech pak pouze 0,27% atd. (Overgaard Knudsen a kol. (2017), str 2). Toto je zvláště omezující na méně rozvinutých trzích, kde je okruh obchodovaných společností (transakcí s podniky) menší ve srovnání s rozvinutými trhy, jak dokázal Nel a kol. (2014) na příkladu Jihoafrické republiky.

Bhojraj a Lee (2002) na tento problém reagovali vytvořením nového metodického konceptu výběru porovnatelné skupiny, který nazývají tzv. warranted multiple (v překladu zaručený násobitel). Bhojraj a Lee (2002) vypočetli sérii ročních regresních analýz s násobiteli EV/S a P/B jako vysvětlované proměnné a osmi fundamenty jako vysvětlující proměnné (následně tuto metodu rozšířili ve verzi An a kol. z roku 2010). Výsledky této

regrese ukazují, že generátory hodnoty (růst, ziskovost a riziko) významně zvyšují sílu modelu ve srovnání s modelem založeným pouze na odvětví (Bhojraj a Lee, 2002, str. 422). Na základě regresních koeficientů z předchozího roku vypočítali warranted multiple který se dá použít přímo pro ocenění anebo na jeho základě identifikovat skupinu nejpodobnějších společností (podle nejpodobnějšího warranted multiple). Bhojraj a Lee (2002, str. 409) udali, že jejich metoda vedla více než k zdvojnásobení upraveného koeficientu determinace (dále též R^2) ve srovnání s PG vybranou na základě odvětví, nebo odvětví kombinovaného s velikostí (u EV/S zvýšení z 18 % na 43 %; u P/B z 4,5 % na 11,3 %) Bhojraj a Lee (2002, str. 427). U společností z tzv. nové ekonomiky (angl. new economy) je zvýšení ještě významnější (v případě EV/S se jedná o zvýšení R^2 o dalších 5-10 %) Bhojraj a Lee (2002, str. 432).

Další metoda, která je založena na podobnosti zvolených fundamentů, je popsána Overgaard Kundsénem a kol. (2017), V ní se pracuje s kritériem nejmenší sumy rozdílů v absolutním pořadí (SARD, z angl. Sum of Absolute Rank Differences) na jednom až pěti fundamentech (RoE, Net Debt/EBIT, velikost tržeb, implikovaný růst a EBIT marže) testovaných na čtyřech různých násobitelích (P/E, P/B, EV/S a EV/EBIT). Overgaard Kundsén a kol. (2017) testovali svou metodu vůči metodě warranted multiple (dle Bhojraj a Lee, 2002) a zjistili superioritu, což vysvětluje menší citlivostí na odlehlá pozorování než regresní analýza, která tvoří základ metody warranted multiple.

Nedávno publikoval svůj článek zaměřený na fundamenty důležité pro výběr PG pro ocenění tržním porovnáním prof. Janda (2019). Konkrétně se zaměřil na jeden z faktorů rizika měřený jako směrodatná odchylka z hospodářských výsledků dané společnosti za 5 let. Na globálních datech za uplynulých 35 let Janda (2019) prokázal, že přidání faktoru stability zisku u vyspělých společností ke kritériím odvětví a geografické lokace zlepšuje odhad P/E a P/B násobitelů vypočtený na základě takto vybrané PG (Janda, 2019, str. 67).

Tabulka 2 shrnuje uvažované proměnné v uvedených článcích.

Tabulka 2 Faktory pro výběr PG podle současné literatury

Článek	Zkoumaní násobitelé	Fundamenty
Alford, 1992	P/E	- odvětví (4ciferný SIC, min však n=6) - bilanční suma - RoE
Cheng a McNamara, 2000	P/E, P/B a a kombinace P/E+P/B	- odvětví (4ciferný SIC, min však n=6) - RoE - bilanční suma
Bhojraj a Lee 2002 (kurzívou pozdější rozšíření: An a kol. 2010)	P/E, P/B (P/E, P/B, EV/S)	- odvětvový násobitel EV/S (<i>odvětvový násobitel podle zkoumaného násobitele</i>) - odvětvový násobitel P/B (----) ** - provozní marže upravená o odvětví * (<i>marže na provoz. zisku</i>) - dummy ztrátovost provozní marže (stejně) - prognózovaný růst - upravený o odvětví * (<i>prognózovaný růst</i>) - zadluženost (<i>stejně</i>) - rentabilita čistých aktiv (<i>RoA</i>) - RoE (<i>stejně</i>) - poměr R&D na výnosech (<i>stejně</i>)
Overgaard Knudsen, 2017	EV/EBIT, EV/S, P/B, P/E	- RoE - Net Debt/EBIT - tržní kapitalizace - prognózovaný růst - EBIT marže - odvětví (6ciferný GICS min však n=16)
Janda, 2019	P/E, P/B	- země - odvětví (3ciferný SIC) - stabilita zisku

(*) Pozn.: fundament „upravený o odvětví“ znamená rozdíl v hodnotě fundamentu mezi daným podnikem a mediánem odvětví.

(**) Pozn.2: v článku ve verzi z roku 2002 byl i ve variantě EV/S zahrnut podíl P/B jako vysvětlující proměnná (kvůli provázanosti P/B na náklady kapitálu dle Gebhardt a kol. 2001 v: Bhojraj a Lee 2002, str. 416). Ve verzi z roku 2010 už byl vždy jen jeden násobitel.

Zdroj: vlastní zpracování

Alternativní metody

Alternativní přístupy jsou různorodé techniky vycházející z toho, že ekonomicky porovnatelné podniky může spojovat nefinanční parametr, který je možné pomocí moderních analytických nástrojů (vč. algoritmů strojového učení) rozklíčovat. Reagují tedy na nedostatky zavedených klasifikačních systémů, které jsem uvedla výše, v podkapitole Agregované skupiny podle odvětvové klasifikace. Protože sestavení porovnatelné skupiny není klíčovým krokem jen pro oceňování (benchmarking se používá často i v dalších oblastech financí, vč. finanční analýzy, auditu, převodních cen a dalších), níže uvedené přístupy se tedy nesnaží sestavit PG výlučně pro valuační účely.

Alternativní metody rozdělení společností do odvětví (stanovení PG) lze ilustrativně rozdělit na: (a) vazby v hledání finančních analytiků, (b) textová analýza a (c) ostatní metody.

a) Vazby v hledání finančních analytiků

Finanční analytici by měli mít díky své znalosti charakteristik jednotlivých společností, které sledují, obecně nejlepší předpoklady k tomu, aby sestavili fundamentálně správnou porovnatelnou skupinu.

Lee a kol. (2015) představili metodiku výběru PG založenou na big data o posloupnosti v internetových vyhledáváních společností v systému EDGAR (Electronic Data-Gathering, Analysis and Retrieval vedeném Komisí pro kontrolu cenných papírů Spojených států amerických (Security Exchange Commission, dále též SEC)). Lee a kol. (2015) vytvořili metodu nazvanou Search Based Peers (dále též SBP) na datech 135,5 mil. webových zobrazení v období od 2008 - 2011 (Lee a kol, 2015). Metoda SBP dominovala výsledkům porovnatelných skupin sestavených na základě 6ciferné klasifikace GICS v šesti různých aspektech, včetně oceňovacích násobitelů (Lee a kol, 2015).

Analýzou PG sestavených v reportech finančních sell-side analytiků se také zabývali De Franco a kol. (2015), kteří však poukázali na to, že analytické reporty tendenčně volí PG s vyššími násobiteli (De Franco a kol., 2015, str. 95). De Franco a kol. předpokládají, že se finanční analytici snaží podpořit investory v nákupu sestavením porovnatelné skupiny s celkově s vyšším oceněním než by odpovídalo fundamentálním charakteristikám oceňovaného podniku (De Franco a kol., 2015, str. 96).

b) Textová analýza

Další uvažovaný spojující nefinanční parametr je popis společnosti (angl. business description). Hoberg a Phillips (2010, a 2016) analyzovali podstatná jména a vlastní jména použitá v produktových popisech v tzv. 10-K závěrkách pomocí nástrojů textové analýzy k nalezení co nejpodobnějších porovnatelných společností. Na svých stránkách Hoberg a Phillips publikují výsledky za období 1996 - 2017, kde pro každou společnost uvádějí individuálně stanovené konkurenty na základě skóre produktové podobnosti (tzv. Text-based Network Industry Classifications).

Podobně Fang a kol. (2013) navrhli přístup, v němž založili klasifikaci odvětví na nalezení slovních podobností v širším popisu činnosti podniku v 10-K závěrkách (tedy nejen popis produktu) s využitím pokročilé metody modelování témat v textu (tzv. Latent Dirichlet Allocation patří do oblasti strojového učení).

S nástroji textové analýzy z oblasti strojového učení pracuje rovněž Chong a Zhu (2012), kteří zpracovávají názvy XBRL značek² (tzv. tags) v účetních závěrkách pro rozlišení různých odvětví. Tento směr dále rozvíjí např. Yang a kol. (2019) s pokročilejšími technikami, pomocí kterých je modelována struktura rozvahy a je poměřována rozdílnost mezi podniky na základě grafického zobrazení struktury.

c) Ostatní

Mezi publikovanými články se objevují i další alternativní přístupy, které jsou v některých případech spíše rarita, např. Xu a kol. (2019), kteří svou odvětvovou klasifikaci založili na mobilitě lidských zdrojů mezi spřízněnými podniky, kterou odvodili z big data sesbíraných na profesních sociálních sítích.

² XBRL (eXtensible Business Reporting Language) je značkovací jazyk, který se prosazuje ve finančním reportingu. Od roku 2008 to je povinný formát pro společnosti vedené u SEC. V roce 2018 schválila Evropská komise plán zavádění pro evropské obchodované společnosti. Detail k vývoji XBRL např. Staňková (2019).

Naopak za velice inspirativní považují nedávno publikovaný článek Ding a kol. (2019) o výběru porovnatelné skupiny pro detekci finančních anomálií založeném na algoritmu strojového učení maximalizující podobnost finančních fundamentů v rámci skupiny. Ding a kol. konkrétně použili metodu shlukování na základě mediánové hodnoty (angl. k-medians clustering method) různých finančních ukazatelů sestavených specificky pro účely detekce finančních anomálií. Metoda shlukování rozdělí všechna pozorování do shluků (angl. clusters) tak, aby vnitroskupinová podobnost byla co nejvyšší. Takto sestavenou porovnatelnou skupinu pak Ding a kol. testovali na dvou situacích: (i) bankrotního predikčního modelu a (ii) modelu detekce chyb (angl. misstatement detection model). V obou výzkumných oblastech dosáhla metoda k-medians clustering zlepšení oproti porovnatelným skupinám sestaveným na základě odvětví a velikosti. Autoři dále doporučují aplikovat k-median clustering i pro další využití, neboť umožňuje dynamicky sestavovat porovnatelnou skupinu na základě kritérií relevantních pro danou situaci. Na tento článek bych chtěla navázat ve smyslu identifikace porovnatelné skupiny s využitím metod strojového učení na základě finančních ukazatelů vybraných specificky pro oceňovací násobitele. Výhodou strojového učení je, že algoritmus díky schopnosti zlepšovat se ze vstupních dat má lepší předpoklady k predikci správné porovnatelné skupiny na základě nelineárních a vzájemně propojených vztahů, které se ve financích často vyskytují.

Strojové učení

S vývojem výpočetní kapacity a rychlého šíření velkých dat (big data) se vyvinuly metody strojového učení (dále rovněž ML, z angl. machine learning). ML je skupina matematických algoritmů na pomezí umělé inteligence a statistických modelů, které „ve strukturovaných a nestrukturovaných (velkých) datech dokáží rozeznat vzájemné vztahy“ (Sellhorn 2020, str. 3). Strojové učení se tedy do určité míry překrývá s klasickými statistickými modely (hranice se nedá přesně určit), ale některé rysy se liší. Hlavní rozdíl spočívá v tom, že místo a priori nastavených předpokladů založených na teoretických předpokladech se ML model z poskytnutých dat sám učí (Ding a kol. 2019b, str. 5). Na jedné straně to vede k výrazně lepší predikční síle ML modelů - díky tomu jsou ML modely někdy označovány jako „prediction machines“ (Sellhorn 2020, str. 14). Na druhou stranu se tím ale ML modely stávají méně transparentní. V závislosti na použité metodě může být z „černé skříňky“ ML algoritmu obtížné pro člověka pochopit, jak ze vstupních dat vznikly výstupní hodnoty (predikce). (Ding a kol. 2019, str. 5, Sellhorn 2020, str. 15-17).

Algoritmy strojového učení se rozlišují na dvě základní skupiny (i) strojové učení s učitelem (angl. supervised ML) a (ii) strojové učení bez učitele (angl. unsupervised ML). Hlavní rozdíl spočívá v tom, zda-li jsou k učení poskytnuta pozorování vysvětlované proměnné.

- a) ML algoritmy učení s učitelem trénují model na základě poskytnutých párů vysvětlované a vysvětlující proměnné. Úkoly pro modely s učitelem jsou především predikce (tzv. regresní úlohy³) a klasifikace (tzv. klasifikační úlohy³) na nových datech (data, která při trénování nebyla použita). Srv. Müller a Guido 2016, str. 2.
- b) Pro ML algoritmy učení bez učitele nejsou k dispozici páry pozorování s vysvětlovanou proměnou, a proto se hodí i pro jiný typ úkolů. Algoritmy bez učitele se obecně využívají k rozpoznání kategorií a popisu souvislostí s cílem snížit dimenzionalitu dat. Srv. Müller a Guido 2016, str. 3. V tomto duchu tzv. shlukovací algoritmus (anglicky clustering) vstupní data rozřadí do skupin (angl. clusters), které se od sebe nějak odlišují, např. bloky textu zařadí do tematických kategorií. Srv. Sellhorn 2020, str. 3.

Strojové učení má široké uplatnění napříč různými obory, včetně financí, kde se s jeho využitím začalo v osmdesátých letech v různých situacích od analýz finančních reportů až po auditní procesy (dle Ding a kol. 2019b, str. 6). Přehled současného výzkumu aplikace ML v oblasti účetnictví a financí podniku detailně zpracoval Sellhorn 2020 a rozdělil je do těchto skupin: (i) účetní odhady založené na ML, (ii) stanovení PG, (iii) identifikace potenciálních dalších transformací ve finančním účetnictví v důsledku ML. Pro tento článek je relevantní aplikace strojového učení v rámci identifikace PG. Současné články k tomuto tématu jsem rozepsala v podkapitole Alternativní metody. Nicméně si nejsem vědoma, že by existovala publikovaná studie využívající strojové učení na sestavení porovnatelné skupiny pro stanovení valuačního násobitele.

Do budoucna pak očekávám, že se potenciál a možnosti strojového učení budou ještě zvyšovat především v důsledku těchto trendů:

- a) Rostoucí množství (anebo vyšší frekvence dostupnosti) strukturovaných informací o podnicích – například viz vývoj finančního a regulatorního reportingu díky technologii XBRL. Jak je detailněji rozepsáno v Staňková (2019), rozvoj XBRL již v některých zemích zahrnuje povinné vykazování i pro neobchodované společnosti (např. Itálie), umožňuje hromadně zpracovávat informace z přílohy účetní závěrky a do budoucna může vést k dostupnosti finančních údajů v reálném čase viz koncept „continuous audit“ (Mejzlík, Ištvánfyová 2006).
- b) Konvergence účetních standardů vedoucí k přesnějšímu srovnávání dat mezi různými regiony – viz dílčí úspěchy konvergenčního úsilí mezi IFRS a US GAAP – a dalších specifických podmínek na finančních trzích v důsledku globalizace.
- c) Vyšší míra zapojení ML algoritmů v různých oblastech, což povede k vyššímu pocitu subjektivní důvěry praktiků.

Nástin budoucího výzkumu

Předmětem plánovaného výzkumu bude ověřit, zda lze s využitím algoritmů strojového učení sestavit model výběru porovnatelné skupiny, resp. jejich násobitele, který bude přesnější než skupina sestavená na základě odvětvového zařazení. Jak jsem uvedla výše, nejsem si vědoma žádného výzkumu, který by se tímto předmětem zabýval, ale vzhledem k rostoucí popularitě strojového učení se již objevují podobné výzkumy, např. Ding a kol. (2019) prokázali zlepšení bankrotního modelu a modelu na detekci chyb v účetnictví díky

³ Rozdíl mezi regresní úlohou a klasifikační úlohou spočívá ve formě vysvětlované proměnné y : (i) regrese: je-li y = spojitá proměnná a (ii) klasifikace: je-li y = kategorická proměnná.

sestavení porovnatelné skupiny algoritmem strojového učení. Na tento výzkum bych chtěla navázat aplikací pro tržní ocenění.

Konkrétně plánuji provést analýzu potenciálu využití strojového učení při stanovení porovnatelné skupiny podniků pro praktické účely ocenění metodou tržního porovnání:

- a) nalezení nejvhodnějšího algoritmu strojového učení k výběru PG na základě vybraných finančních ukazatelů potenciálně použitelného v praxi,
- b) otestování tohoto ML modelu oproti modelu primárně založenému na odvětvovém zařazení (reprezentující současnou manuální praxi) a
- c) v případě kladného výsledku pak cílovým výstupem bude specifikace doporučeného modelu, který bude možné v praxi používat jako komplementární úhel pohledu na výběr PG, zároveň předpokládám sestavení výzkumných otázek pro navazující výzkum.

Faktory pro výběr PG

Pro svůj výzkum předpokládám využití následujícího výběru proměnných pro sestavení ML modelu k identifikaci PG pro ocenění tržním porovnáním. Tento výběr jsem sestavila především na základě současné literatury zabývající se tímto tématem shrnuté v sekci 2.1. Výběr porovnatelné skupiny podniků v kontextu ocenění tržním porovnáním a především v Tabulce 2:

(i) RoE, (ii) Net Debt/EBIT, (iii) tržní kapitalizace (proxy k velikosti), (iv) prognózovaný růst tržeb dle databáze IBES normalizovaný dle Overgaard Knudsen 2017 a An a kol. 2010, (v) EBIT marže, (vi) odvětví v různé míře podrobnosti (na základě Overgaard Knudsen 2017), (vii) statutární daňovou sazbu z příjmu právnických osob (CIT) a (viii) bilanční sumu (TA), neboť na základě Alforda (1992) rovněž efektivně nahrazuje rizikovost vyjádřenou betou, která je opět dostupná pouze pro obchodované společnosti) a (ix) obrat (Rev) jako dvě alternativní proxy k velikosti, aby bylo možné model aplikovat i na neobchodované společnosti, (x) dummy proměnná pro ztrátové společnosti (Loss) podle Bhojraj a kol. 2010, (xi) EBITDA marže, (xii) stabilita zisku (dle Janda 2019), (xiii) RoA (alternativa k RoE pro brutto násobitele), (xiv) zadluženost poměrem dluh vůči vlastnímu kapitálu (další proxy k riziku) a (xv) geografický region, (xvi) poměr investičních výdajů k tržbám.

V úvahu připadají ještě další proměnné, ale buď se obávám, že by nebyly dostupné pro mnoho společností, čímž by zmenšovaly počet pozorování, nebo jsou obtížně zpracovatelné pro jednoduchý model (např. náklady na výzkum a vývoj, vlastnická struktura) anebo předchozí studie prokázaly statistickou nesignifikantnost (poměr vyplacených dividend viz Bhojraj a Lee 2002). Zdrojová data předpokládám získat v databázi Thomson Reuters.

Model strojového učení

Výhodou strojového učení (ML) oproti běžné regresní analýze ve statistice je výrazně vyšší flexibilita, která je dána tím, že ML nepotřebuje dopředu stanovit tolik parametrů modelu, protože ML sám hledá vztahy mezi pozorováními.

Ve svém výzkumu předpokládám využití dvou základních přístupů k ML modelování: (i) shlukování (angl. clustering) v rámci skupiny algoritmů „strojového učení bez učitele“ a (ii) regrese v rámci kategorie „strojového učení s učitelem“. Pomocí těchto modelů lze nezávisle na sobě odhadnout násobitele $\widehat{m}_i$.

Testování kvality / přesnosti modelu v rámci strojového učení probíhá na základě porovnávání s výsledky na kontrolním vzorku dat. Část dat se ponechá stranou a v závěrečném hodnocení modelu se porovnává skutečná hodnota násobitele m_i s vypočtenou (odhadnutou) implikovanou hodnotou násobitele $\widehat{m}_i$ na základě finančních parametrů (toto u všech testovaných násobitelů). Výslednou přesnost modelu pak vyčíslím jako absolutní procentuální chybu podle následujícího vzorce.

$$P = \frac{1}{n} \sum_{i=1}^n |\widehat{m}_i - m_i| \quad (1)$$

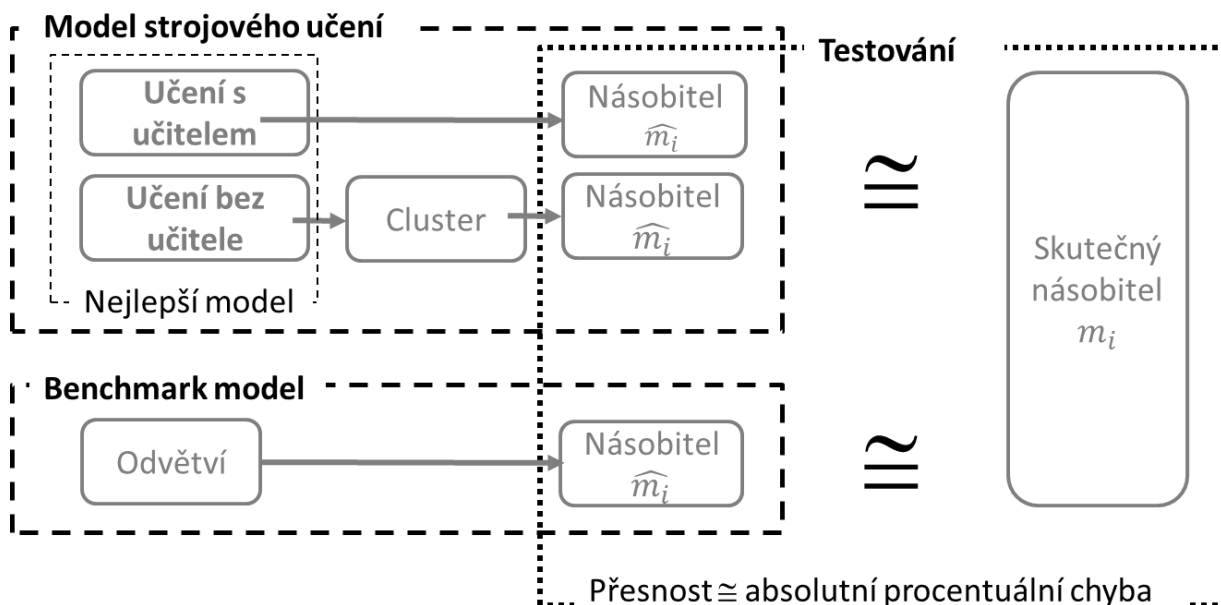
kde P – přesnost,
 $\widehat{m}_i$ – odhadnutý násobitel společnosti i ,
 m_i – skutečný násobitel společnosti i ,

Benchmark model

Efektivnost modelu strojového učení lze prokázat nejenom na základě přesnosti odhadů ve srovnání se skutečností, ale především na srovnání s alternativní metodou. Jak uvádím v teoretické části, v současnosti převládá metoda založená na zařazení do odvětví, proto jsem si jako benchmark model zvolila přístup Agregované skupiny podle odvětvové klasifikace. Pomocí tohoto přístupu rovněž vypočítám (odhadnutou) hodnotu násobitele $\widehat{m}_i$, který porovnám se skutečnou hodnotou násobitele m_i , abych zjistila přesnost podle rovnice (1). Na základě porovnání přesnosti (i) modelu strojového učení a (ii) benchmark modelu se pak budu moci vyjádřit k potenciálu využití strojového učení při výběru porovnatelné skupiny podniků pro tržní ocenění.

Předpokládaný design modelu schématicky znázorňuje Obrázek 1.

Obrázek 1 Předpokládaný design modelu



Zdroj: vlastní zpracování

Závěr

V tomto článku jsem se zabývala problematikou výběru porovnatelné skupiny podniků. V rámci tržního ocenění se jedná o zásadní krok, na němž v podstatě závisí kvalita výsledného ocenění. Cílem tohoto článku bylo popsat současnou literaturu, která nabízí různé přístupy k výběru srovnatelných podniků. V praxi nejčastěji převládá postup založený na odvětvové klasifikaci, a to i přes určité nevýhody, které tomuto přístupu teorie vytýká. V rámci připravovaného výzkumu plánuji otestovat tento v praxi nejběžnější postup vůči modelu využívající strojové učení. Budu ověřovat, jestli metody strojového učení, které jsou založené na schopnosti predikovat na základě pozorovaných vztahů v poskytnutých vstupních datech namísto a priori nastavených předpokladech, mají potenciál do budoucna zpřesnit výběr porovnatelné skupiny.

Literatura

- [1] Alford, A.W. (1992): *The Effect of the Set of Comparable Firms on the Accuracy of the Price-Earnings Valuation Method*. Journal of Accounting Research 30, 94–108. <https://doi.org/10.2307/2491093>
- [2] An, J., Bhojraj, S., Ng, D.T. (2010): *Warranted Multiples and Future Returns*. Journal of Accounting, Auditing & Finance 25, 143–169. <https://doi.org/10.1177/0148558X1002500201>
- [3] Bhojraj, S., Lee, C.M.C., (2002): *Who Is My Peer? A Valuation-Based Approach to the Selection of Comparable Firms*. J Accounting Res 40, 407–439. <https://doi.org/10.1111/1475-679X.00054>
- [4] Bhojraj, S., Lee, C.M.C., Oler, D.K., (2003): *What's My Line? A Comparison of Industry Classification Schemes for Capital Market Research*. J Accounting Res 41, 745–774. <https://doi.org/10.1046/j.1475-679X.2003.00122.x>
- [5] Cheng, C.S.A., McNamara, R., (2000): *The valuation accuracy of the Price/Earnings and Price-Book Benchmark Valuation methods*. Review of Quantitative Finance and Accounting 15, 349–370. <https://doi.org/10.1023/A:1012050524545>
- [6] Damodaran, A., (2002): *Multiples: First Principles*.
- [7] De Franco, G., Hope, O.-K., Larocque, S., (2015): *Analysts' choice of peer companies*. Rev Account Stud 20, 82–109. <https://doi.org/10.1007/s11142-014-9294-7>
- [8] Ding, K., Peng, X., Wang, Y., (2019): *A Machine Learning-Based Peer Selection Method with Financial Ratios*. Accounting Horizons 33, 75–87. <https://doi.org/10.2308/acch-52454>
- [9] Ding, K., Lev, B.I., Peng, X., Sun, T., Vasarhelyi, M.A., (2019b): *Machine Learning Improves Accounting Estimates*. Presented at the Review of Accounting Studies Conference, Singapore.
- [10] Dittmann, I., Weiner, C., (2005): *Selecting Comparables for the Valuation of European Firms*. SSRN Journal 25 stran. <https://doi.org/10.2139/ssrn.644101>
- [11] Fang, F., Dutta, K., Datta, A., (2013): *LDA-based industry classification*. Presented at the Thirty Fourth International Conference on Information Systems, Milan.
- [12] Hoberg, G., Phillips, G., (2016): *Text-Based Network Industries and Endogenous Product Differentiation*. Journal of Political Economy 124, 1423–1465. <https://doi.org/10.1086/688176>

- [13] Hoberg, G., Phillips, G., (2010): *Product Market Synergies and Competition in Mergers and Acquisitions: A Text-Based Analysis*. Rev. Financ. Stud. 23, 3773–3811. <https://doi.org/10.1093/rfs/hhq053>
- [14] Janda, K., (2019): *Earnings stability and peer company selection for multiple based indirect valuation*. Finance a Uver - Czech Journal of Economics and Finance 69, 37–75.
- [15] Knudsen Overgaard, J., Kold, S., Plenborg, T., (2017): *Stick to the fundamentals and discover your peers*. Financial Analysts Journal 73, 85–105.
- [16] Lee, C.M.C., Ma, P., Wang, C.C.Y., (2015): *Search-based peer firms: Aggregating investor perceptions through internet co-searches*. Journal of Financial Economics 116, 410–431. <https://doi.org/10.1016/j.jfineco.2015.02.003>
- [17] Mařík, M., (2011) *Metody oceňování podniku: proces ocenění - základní metody a postupy*. Ekopress, Praha.
- [18] Mejzlík, L., Ištvánfyová, J., (2006): *Vykazování podnikových dat v elektronické podobě – XBRL*. Český finanční a účetní časopis, 2006.2: p. 124-129. <https://doi.org/10.18267/j.cfuc.157>
- [19] Müller, A.C., Guido, S., (2016): *Introduction to machine learning with Python: a guide for data scientists*, First edition. ed. O'Reilly Media, Inc, Sebastopol, CA.
- [20] Nel, S., Bruwer, W., le Roux, N., (2014): *An emerging market perspective on peer group selection based on valuation fundamentals*. Applied Financial Economics 24, 621–637. <https://doi.org/10.1080/09603107.2014.894629>
- [21] Plenborg, T., Pimentel, R.C., (2016): *Best Practices in Applying Multiples for Valuation Purposes*. JPE 19, 1–10. <https://doi.org/10.3905/jpe.2016.19.3.055>
- [22] Sellhorn, T., (2020): *Machine Learning und empirische Rechnungslegungsforschung: Einige Erkenntnisse und offene Fragen*. Schmalenbachs Z betriebswirtsch Forsch 72, 49–69. <https://doi.org/10.1007/s41471-020-00086-1>
- [23] Staňková, V., (2019): *Electronic data format XBRL: recent development and prospects in the EU*. Presented at the The 20th Annual Conference on Finance and Accounting, Praha.
- [24] Vydržel, K., Soukupová, V., (2012): *Empirical examination of valuation methods used in private equity practice in the Czech Republic*. Journal of Private Equity 16, 83–99. <https://doi.org/10.3905/jpe.2012.16.1.083>
- [25] Xu, X., Qian, H., Ge, C., Lin, Z., (2019): *Industry classification with online resume big data: A design science approach*. Information & Management 103182. <https://doi.org/10.1016/j.im.2019.103182>
- [26] Yang, S.Y., Liu, F., Zhu, X., Yen, D.C., (2019): *A Graph Mining Approach to Identify Financial Reporting Patterns: An Empirical Examination of Industry Classifications*. Decision Sciences 50, 847–876. <https://doi.org/10.1111/deci.12345>

K problému výběru porovnatelné skupiny podniků pro ocenění a nástin budoucího výzkumu s využitím strojového učení

Veronika Staňková – Miloš Mařík

ABSTRAKT

Tento článek se zabývá problematikou výběru porovnatelné skupiny podniků pro ocenění tržním porovnáním. Akademická odborná veřejnost se neshoduje na optimálním přístupu k výběru. V článku proto nejprve uvádíme syntézu literatury k tomuto tématu, včetně popisu a diskuze ke třem hlavním přístupům: (i) agregované skupiny podle odvětvové klasifikace, (ii) hledání fundamentálních ukazatelů a (iii) alternativní metody využívající big data. Dále se zabýváme tématem strojového učení aplikovaného ve financích a přinášíme nástin budoucího výzkumu, ve kterém chceme ověřit potenciál využití strojového učení při výběru porovnatelné skupiny podniků.

Klíčová slova: Ocenění tržním porovnáním; Porovnatelná skupina společností; Strojové učení.

Selecting a peer group of companies for valuation and outline of future research using machine learning

ABSTRACT

This article deals with peer group selection for the purpose of the market valuation approach. The academic professional public does not agree on the optimal approach in respect of the peer group selection. Therefore, in this article we start with a synthesis of the current literature on this matter, including a description and discussion of the three main approaches, namely: (i) aggregated groups based on industry classification, (ii) search for fundamental indicators and (iii) alternative methods using the big data. Following is the explanation of machine learning applied in the finance and an outline of future research in which we want to verify the potential of machine learning algorithms in selecting a comparable group of companies.

Key words: Market valuation approach; Peer group; Machine learning.

JEL classification: G120