

# Strojové učení v diagnostice a při predikci časových řad

Jiří Kléma

Gerstnerova laboratoř pro inteligentní rozhodování a řízení,  
katedra kybernetiky, České vysoké učení technické, elektrotechnická fakulta,  
Technická 2, 166 27 Praha 6, Česká republika  
klema@labe.felk.cvut.cz

**Abstrakt.** Strojové učení (ML - Machine Learning) dosahuje úspěchů v celé řadě praktických aplikací, což z něj činí populární a často využívanou třídu metod. Příspěvek se zabývá nejprve teoretickým pohledem na strojové učení a jeho algoritmy a metody. Obecně zhodnocuje, na jaké úrovni je současná adaptivní schopnost učících se systémů samotných a nakolik je vývoj nové aplikace podporován kooperativním úsilím znalostního inženýra a experta v dané problémové oblasti. V druhé části poté demonstruje některé úspěšné aplikace strojového učení realizované v rámci Gerstnerovy laboratoře. Hlavní pozornost je věnována úloze predikce spotřeby plynu, krátce je diskutována také úloha diagnostiky chybových stavů čerpadla.

## 1 Strojové učení

Strojové učení (ML - Machine Learning) dosahuje úspěchů v celé řadě praktických aplikací, což z něj činí populární a často využívanou třídu metod. Cílem tohoto článku je zhodnotit příčinu těchto úspěchů, ale také naznačit současné limity tohoto oboru a jeho příští vývoj.

### 1.1 Úvod, obecná a konkrétní definice

Z teoretického hlediska existuje celá řada definic pojmu strojového učení. Kohavi a Provost [10] definují strojové učení jako obor, který se zabývá induktivními, ale i jinými algoritmy, které mohou být považovány za učící se. Tato jistě platná definice ovšem ještě nenaznačuje, co vlastně za učení v případě počítačových systémů (programů, algoritmů) považujeme. Právě specifikace pojmu učení totiž nejlépe dokumentuje stav, schopnosti i příští vývoj tohoto oboru. Mitchell [12] definuje pojem učení jako schopnost systému zlepšovat svoji výkonnost na dané třídě úloh se získáváním zkušeností. Tato definice shrnuje hlavní představy o učení. Je založena na obecných pojmech zkušenost, třída úloh a výkonnost. Dřívější vývoj oboru ale prokázal, že přímé naplnění smyslu této obecné definice není nikterak jednoduché. Strojové učení tedy v rámci své teorie často sleduje konkrétnější definici. Tou je snaha o indukci obecného konceptu (funkce) na základě konkrétních trénovacích příkladů. Hovoříme pak o takzvaném učení z příkladů.

### 1.2 Cyklus návrhu aplikace, problémy a moderní přístupy k jejich řešení

Z hlediska praktických úloh se může jednat například o dlouhodobé sledování denních odběrů plynu v distribuční síti, kdy je naším cílem nalézt (naučit se) funkční vztah mezi měřitelnými veličinami (atributy) a veličinou cílovou (predikovanou). Algoritmus strojového učení, ve smyslu své konkrétní definice, dostává na vstupu množinu trénovacích příkladů, výstupem je pak znalost reprezentovaná modelem učené funkce. Model je následně využíván pro predikci spotřeby ve dnech, kdy tato hodnota není známa. Je ovšem zřejmé, že při aplikaci strojového učení je výše uvedený algoritmus třeba doplnit o další kroky. Jedná se zejména o:

- formulaci problému – tj. stanovení učeného konceptu (funkce),
- specifikaci trénovacích příkladů – tj. určení sledovaných nezávislých proměnných (atributů), které charakterizují trénovací příklady, tyto atributy často nejsou přímo měřitelné a mohou být i složitě odvozovány,
- shromažďování zkušeností – tj. vlastní sběr trénovacích příkladů a posuzování jejich relevance v řešené doméně,
- vyhodnocení naučené znalosti – tj. zhodocení kvality získaného modelu, stanovení minimálního počtu trénovacích příkladů apod.; funkce hodnotící kvalitu naučeného modelu velice často významně ovlivňuje i samotný model (tzv. wrapper approach, kdy je model iterativně modifikován s cílem maximalizovat hodnotící funkci),
- aplikace modelu – zasazení odvozené znalosti do uživatelského prostředí, popř. integrace do existujícího znalostního či řídicího systému.

Přes značné úsilí je úspěšné naplnění výše uvedených kroků stále nejčastěji výsledkem kooperativního úsilí znalostního inženýra a experta v dané problémové oblasti. Nejedná se tedy o automatizované postupy. Tím může buď strádat inteligence a adaptivita navržených systémů nebo to přinejmenším činí jejich návrh zdoluhavým a obtížným. Příkladem může být to, že se systém není schopen vyrovnat s nečekanými a neočekávanými změnami prostředí nebo výpadky vstupních údajů a to i tehdy, pokud je inteligentní adaptace v zásadě možná. Úspěšnost systému pak není limitována kvalitou vlastního učícího se algoritmu, ale spíše správností provedení daných kroků. Výše uvedené paradoxy strojového učení jsou podrobně diskutovány v článku Jindřicha Bůchy [4].

Obtížnost univerzální automatizace těchto kroků je zřejmá, jejich důležitost také. Umělá inteligence se pokouší zmíněné paradoxy řešit na dvou úrovních. V první úrovni se strojové učení stává součástí širšího postupu, který formalizuje celý proces návrhu nové aplikace. Dochází tedy k jeho zařazení do nebo prolnutí s dobýváním znalostí (DM - Data Mining). Tato problematika je studována v celé řadě projektů výzkumných projektů (mj. evropský projekt Data Mining and Decision Support for business competitiveness: Solomon European virtual enterprise (Sol-Eu-Net)).

Druhou úroveň řešení je snaha o významné rozšíření schopností učících se systémů jako takových. V poslední době se proto pozornost soustřeďuje na výzkum inteligentních adaptivních systémů (SAS – Smart Adaptive Systems), které splňují požadavky obecné definice strojového učení. Jsou zároveň schopny převzít alespoň část úkolů vyhrazených dosud výhradně inteligenci lidské. Významná pozornost je jim věnována i v rámci EUNITE Network [1]. Za inteligentní a adaptivní není považován systém, u něhož proběhne učení jednorázově a systém se již následně nevyvíjí. Za takovýto systém tedy nelze považovat například neuronovou síť, jejíž struktura a váhy jsou již po předchozím učení neměnné. Klíčovým pojmem je u těchto systémů schopnost evoluce – tedy procesu modifikace v čase a v reakci na okolní prostředí. V rámci pracovní definice SAS se rozlišují tři úrovně adaptace [1]:

- přizpůsobení se drobné změně prostředí – nejjednodušší úroveň adaptace, která může být postižena přeučněním modelu reprezentujícího cílovou funkci; důraz je kladen na automatickou evokaci přeučení a samostatnou volbu nového modelu; může jí být dosaženo například rozlišením krátkodobé a dlouhodobé paměti systému v kombinaci se spolehlivou metodou hodnotící výkonnost momentálního nastavení modelu,
- přizpůsobení se podobnému prostředí bez explicitního upozornění na tuto změnu – nejde již o pouhý posun významu jednotlivých atributů, ale důraz je kladen na změnu prostředí samotnou; může se jednat například o
- přizpůsobení se nové/dříve neřešené aplikaci – v současné době futuristická vize, systém lze automaticky vytvořit na základě omezeného množství apriorní znalosti s využitím inkrementálního učení; systém se sám vyvíjí a autonomně hledá řešení dříve neřešeného problému.

Pokud shrneme současnou úroveň většiny aplikací strojového učení ve smyslu výše uvedených argumentů, jedná se zpravidla o systémy neadaptivní či s nízkou úrovní adaptace. To v žádném případě neznamená neúspěšnost konečného řešení. Tento text si také nedává za cíl odrazovat od aplikací strojového učení, ale naopak přiblížit strojové učení a jeho současné a příští možnosti průmyslovému uživateli.

Následující kapitoly demonstrují dvě úspěšné aplikace strojového učení realizované v rámci Gerstnerovy laboratoře. Konkrétně se jedná o úlohy diagnostiky chybových stavů čerpadla a predikce spotřeby plynu. Jedná se o úlohy řešené tradičním způsobem, konečný výsledek je tedy založen nejen na kvalitě zvoleného učícího se algoritmu, ale také kvalitou analýzy řešené úlohy, předzpracování dat a vhodnosti zvoleného algoritmu pro řešenou úlohu.

## 2 Predikce spotřeby plynu

### 2.1 Motivace, základní algoritmy

Lokální (místní) distributoři plynu zpravidla odebírají plyn od národních dodavatelů a dále jej distribuují v rámci regionu, který pokrývají. Jejich kontrakt s národními dodavateli je obvykle limitován a určen smluvními denními limity odběru. Lokální distributoři musí každodenně udržovat své odběry v těchto předem dohodnutých mezích. Kdykoli je limit překročen, lokální distributor musí platit značně vysokou pokutu. Je zřejmé, že denní časová predikce množství odebraného plynu může lokálním distributorům značně napomoci regulovat svůj odběr v rámci nasmlovaných mezí. Zpravidla mají dvě možnosti, jak se penalizaci vyhnout. Jednak mohou některým svým odběratelům po dohodě doporučit přepnutí na alternativní zdroj energie (např. olej). Druhá možnost, jak se vyhnout penalizaci, je naplnit lokální rezervní nádrže plynu během období menších spotřeb. Zásoby z nádrží se potom použijí během dnů se zvýšeným odběrem plynu [9].

Úloha predikce spotřeby plynu nepochybně patří do kategorie úloh, kde se bez využití předchozích zkušeností a tím i učení, uspokojivého řešení dosáhnout nedá. Při přesnější kategorizaci je predikce spotřeby plynu úlohou numerické predikce časové řady. V technické praxi lze nalézt celou řadu analogických úloh – předpověď počasí (teploty, sluneční aktivity atd.), míry nezaměstnanosti, směnných kurzů valut, ceny akcií – zejména poslední dvě finanční aplikace jsou předmětem eminentního zájmu komunity umělé inteligence, a proto lze metody jejich řešení nalézt v celé řadě studií. Důležitou charakteristikou tohoto druhu úloh je ovlivnění koncového výstupu (v našem případě denního odběru plynu) celou řadou různých faktorů, z nichž některé nemusí být explicitně zmíněny při popisu úlohy. Zároveň neexistuje žádná jednoznačná analytická funkční závislost mezi minulými a současnými hodnotami uvedené časové posloupnosti ani podobná závislost mezi současnou hodnotou predikované veličiny a ovlivňujícími faktory. Strojové učení nabízí pro tyto úlohy celou řadu teoreticky velice důkladně propracovaných metod – případové usuzování (CBR – case-based reasoning), neuronové sítě (NN – neural nets) a metoda k-nejbližších sousedů (kNN – k-nearest neighbours) patří mezi dlouhodobě nejúspěšnější a nejvyhledávanější. Současně se osvědčila kombinace metod umělé inteligence se statistickými regresními metodami extrapolace.

Metoda CBR je jednou ze základních metod automatického uvažování a strojového učení. Její podstata spočívá v tom, že řešitel přistupuje k problému tak, že se snaží nalézt jeho podobnost s problémy, jejichž řešení je již známo. Nehledá tak pro nový problém nová unikátní řešení, ale snaží se jej odvodit z modifikací řešení úspěšných v minulosti (z minulých příkladů). Uplatníme-li metodu CBR na systém predikce plynu, tak za nový problém považujeme nový den připravený k predikci, za vyřešené problémy se pokládají minulé dny s již známými denními odběry plynu. Jednotlivé dny jsou uloženy jako příklady, které reprezentují specifickou znalost a vytvářejí model systému.

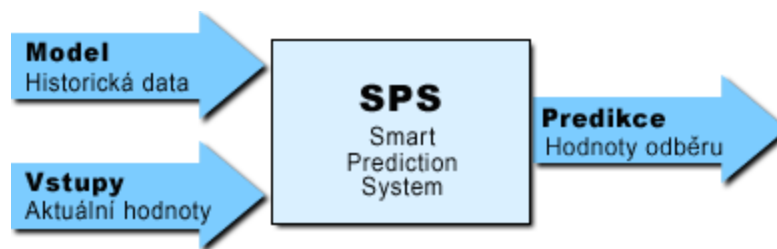
Neuronové sítě patří také mezi populární modely pro predikci časových řad. Nejpoužívanější modifikací neuronových sítí je vícevrstvá síť perceptronů. Díky univerzálním aproximačním schopnostem jsou neuronové sítě schopné predikovat (aproximovat) i velmi složité vztahy mezi vstupy a výstupy. Bylo experimentálně ověřeno, že jsou také schopné řešit nelineární vztahy mezi vstupy a výstupy.

Většina modulů využívajících statistického zpracování dat využívá lineární regresní analýzy pomocí rozšířené metody nejmenších čtverců. Tato metoda hledá nejideálnější lineární kombinace atributů (vstupů) vzhledem k jednotlivým výstupům. Z analýzy dat často může vyplynout potřeba zavést do systému uměle ne-

linearitu, která by lépe reflektovala charakteristiky systému. Takto obohacený nelineární systém dosahuje lepších výsledků a predikční přesnosti, a proto se také využívá v případě predikce odběru plynu. Nelinearita je zaváděna pomocí transformačních funkcí aplikovatelných na jednotlivé vstupní atributy.

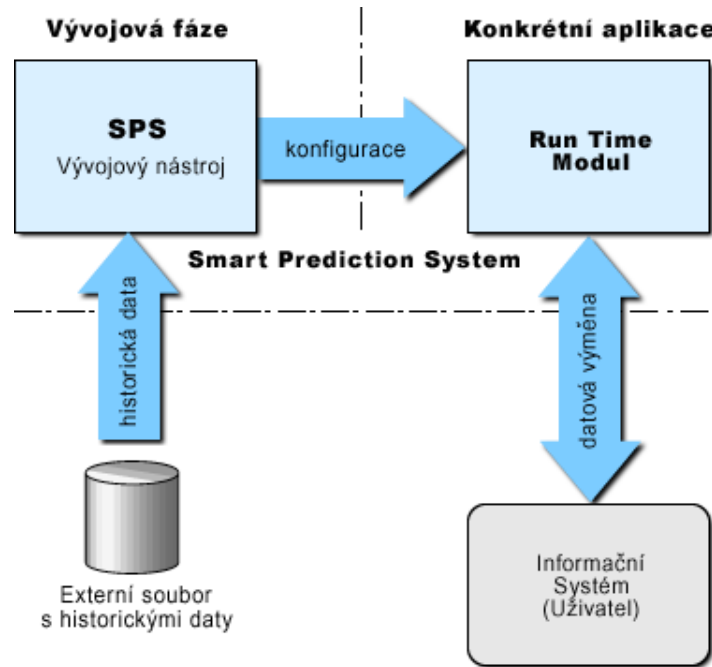
## 2.2 Navržené průmyslové řešení – systém SPS

Problematika predikce spotřeby plynu je průmyslově řešena v rámci systému SPS (Smart Prediction System). Na základě historických dat o chování dané distribuční sítě v minulých dnech (letech) se v rámci SPS systému vytvoří prediktivní model. Ten pak generuje predikce spotřeby plynu na základě aktuálních hodnot veličin měřených v dané distribuční síti (viz Obr. 1). Na druhou stranu, přes výše zmíněnou jednoduchost může být úloha predikce definována v celé řadě modifikací. Modifikace se mohou promítnout např. do délky období, pro které má být spotřeba predikována (hodina, den, týden), do počtu a struktury zaznamenávaných ovlivňujících faktorů (teplota, rychlost větru, vlhkost vzduchu, oblačnost nebo regulační zásahy obsluhy). Výsledné řešení může být také významně ovlivněno délkou období, po které jsou spotřeba a sledované ovlivňující faktory zaznamenávány. Tato délka by měla být tím delší, čím je delší období, pro které má být spotřeba predikována a to alespoň v řádu stonásobků.



Obr. 1. Základní pohled na SPS

V nejčastější podobě zadání distributoři plynu požadují 24 hodinovou predikci s tím, že je tato predikce v průběhu 24 hodin dále zpřesňována s nově se objevujícími aktuálními údaji o spotřebě. Vedle časové řady minulých odběrů plynu je jako ovlivňující faktor zpravidla zaznamenávána venkovní teplota, jež ze škály ovlivňujících faktorů vykazuje statisticky nejvyšší korelaci se spotřebou plynu. Míra této korelace pro různé distribuční sítě kolísá, vliv teploty na spotřebu je však evidentní u všech typů sítí. SPS se zaměřuje právě na výše zmíněný typ zadání. Základní časovou jednotkou je tedy z hlediska predikce jeden plynový den (od 6am do 6am následujícího dne), SPS predikuje 24 hodinovou spotřebu. Každý den je charakterizován 24 hodnotami hodinových odběrů a 24 hodnotami venkovní teploty.



**Obr. 2.** Vývoj nové aplikace predikčního systému – struktura SPS

Z výše uvedeného popisu plyne, že prediktivní systém pracuje ve dvou základních režimech – učení a vlastní predikce. V rámci učení se vytváří prediktivní model odpovídající dané distribuční síti, režim predikce pak tento model aplikuje na aktuální data a generuje odhady příští spotřeby. V rámci systému SPS hovoříme o vývojovém nástroji a run-time modulu. Vývojový nástroj je používán ve fázi učení na jejímž závěru je vygenerován run-time modul, který slouží k vlastní predikci. Struktura systému SPS je naznačena na Obr. 2.

SPS vývojový nástroj má bohaté uživatelské rozhraní, které umožňuje načtení historických dat, dále volbu prediktivních algoritmů použitých v daném konkrétním případě a jejich následné trénování. Na závěr je definován tzv. meta-prediktor, který na základě jednoduchých pravidel kombinuje výsledky jednotlivých dílčích algoritmů. Všechny výše uvedené kroky jsou podporovány dialogy, samozřejmostí je i vizualizace průběžných i konečných výsledků. Práce s vývojovým nástrojem přesto vyžaduje jistou znalost jednotlivých prediktivních algoritmů, předpokládá se tedy, že s vývojovým nástrojem pracuje expert (zaškolený pracovník).

Run-time modul je programový kód, který je exportován z vývojového nástroje po ukončení tréninku a ověření správnosti dosahovaných odhadů spotřeby. U tohoto modulu se předpokládá integrace do informačního (řídícího) systému zákazníka, modul tedy neobsahuje uživatelské rozhraní. Na vstup modulu přicházejí z informačního systému aktuální hodinové údaje o spotřebě a venkovní teplotě, na začátku dne také odhad teploty na příštích 24 hodin, výstupem je po hodinách upřesňovaná predikce denní spotřeby plynu v dané distribuční síti. Pro využití run-time modulu není třeba žádná speciální znalost, je nutné pouze zajistit správnost vstupujících údajů.

## 2.3 Vývoj nové aplikace

Vývoj nové prediktivní aplikace pro novou distribuční síť může být rozložen na následující kroky:

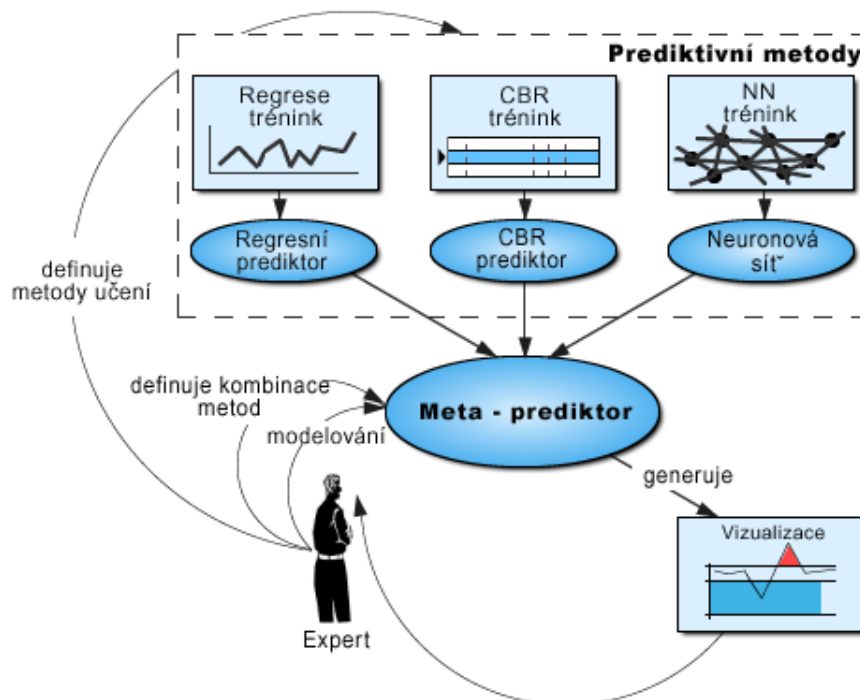
1. sběr dat (pokud údaje za uplynulé roky nebyly uchovávány),
2. předzpracování historických dat do formátu externího souboru,
3. předzpracování dat v rámci vývojového nástroje SPS – definice vstupů prediktivních algoritmů,
4. trénink individuálních prediktivních algoritmů,

5. definice a testování optimálního meta-prediktoru,
6. verifikace kvality predikce,
7. export run-time modulu,
8. integrace run-time modulu do informačního systému zákazníka,
9. vlastní on-line predikce (použití vytvořené aplikace).

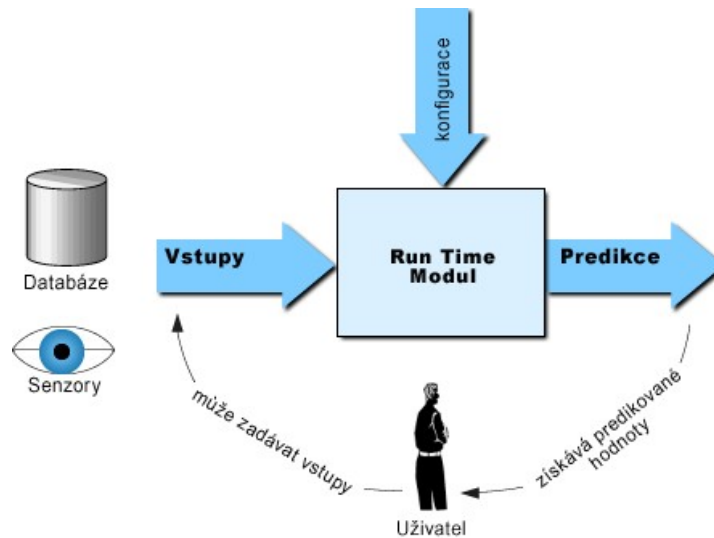
Kroky 3 až 7 jsou realizovány v rámci vývojového nástroje (viz. Obr. 3). Vývojové prostředí SPS pracuje na svém vstupu s tzv. archivním souborem, který obsahuje údaje o chování distribuční sítě v uplynulém období. Soubor je vytvářen mimo SPS (kroky 1 a 2). Zkušenosti s dřívějším použitím SPS naznačují, že pro většinu sítí je optimální a dostačující tříleté monitorovací období. Relativně dobrých výsledků lze dosáhnout i na základě archivního souboru obsahujícího data za jeden rok. Použití databází a uchovávání informací je dnes standardním postupem, archivní soubor tedy zpravidla může být vytvořen okamžitě bez dlouhodobého sběru dat.

Archivní soubor je členěn do jednotek, tzv. trénovacích příkladů. Struktura trénovacích příkladů odpovídá délce predikovaného období. SPS generuje denní predikce, jeden trénovací příklad tedy reprezentuje také 24 hodinovou periodu. Každý trénovací příklad se dále dělí na jednotlivé atributy (den, měsíc, rok, průměrná denní teplota, průměrná denní teplota za minulý den, ...). Struktura trénovacího příkladu je jednoznačně dána, skutečné využití jednotlivých atributů pak může být upraveno v rámci vývojového prostředí SPS.

Po naučení jednotlivých algoritmů a vytvoření meta-prediktoru je dalším krokem export run-time modulu. Jednotlivé run-time moduly mají stejné jádro ale odlišují se konfigurací. Součástí konfigurace je jednak archivní soubor a dále parametry určující chování jednotlivých metod i celého meta-prediktoru. Důležitou složkou je z hlediska spolehlivosti využití predikčního systému jeho správná integrace do celkového informačního (řídícího) systému zákazníka. Integrace musí zajistit především budoucí správnost dat vstupujících do predikčního systému (run-time modulu) a také integritu při aktualizaci archivního souboru (například v souvislosti s opakovaným spuštěním). Použití run-time modulu je schematicky naznačeno na Obr 4.



**Obr. 3.** Vývoj nové aplikace predikčního systému – struktura SPS



**Obr. 4.** Run-time modul – konkrétní aplikace predikčního systému

## 2.4 Shrnutí

Systém SPS byl vyvinut pro lokální dodavatele plynu v Německu ve spolupráci s německou firmou TeleData-Electronics. V současné době je úspěšně využíván čtyřmi lokálními dodavateli: Paderborn, NGW, Rheide, Buckeburg (celkem šest nezávislých distribučních sítí). Úspěšnost systému závisí na charakteru distribuční sítě. Na základě experimentálních zkušeností bylo ověřeno, že průměrná chyba predikce se na začátku predikovaného dne pohybuje v jednotkách procent a v průběhu dne dále klesá. Některé sítě vykazují silný vztah mezi klimatickými podmínkami a skutečnou spotřebou, jiné sítě vykazují vlivem odlišné struktury odběratelů větší nahodilost skutečné denní spotřeby plynu. Důležitým faktem ale je, že vždy byly splněny požadavky zákazníka, který systém SPS využívá k řízení své distribuční sítě.

Z teoretického hlediska se jedná o problémově orientovaný integrovaný systém. SPS je moderním nástrojem, který integruje nejvhodnější algoritmy numerické predikce s ohledem na specifika úlohy předpovědi spotřeby plynu (vztah hodinových a průměrných hodnot, postupné zpřesňování predikce). Dále využívá technologie meta-predikce, která využívá nejnovějších poznatků z oblasti strojového učení. Zjednodušeně lze uvést, že kombinace několika různých algoritmů poskytuje spolehlivější a přesnější výsledky než dílčí algoritmy. Z hlediska návrhu nové aplikace sice nevykazuje automatickou adaptivitu na změnu prostředí, ale nabízí efektivní uživatelské prostředí integrující všechny kroky nutné pro vývoj příbuzné aplikace.

Z technického hlediska SPS zahrnuje uživatelsky příjemné prostředí pro vývoj nových aplikací a zároveň umožňuje automatické generování efektivních run-time modulů. Tyto moduly jsou určeny pro spolupráci se stávajícími systémy zákazníka. Stávající aplikace SPS potvrzují, že včasný a správný odhad spotřeby plynu v dané distribuční síti následovaný odpovídajícím regulačním opatřením, vede ke značným finančním úsporám daného distributora plynu.

### 3 Diagnostika čerpadla

V této úloze bylo cílem vytvořit systém pro včasnou neinvazivní diagnostiku chybových stavů čerpadla. Celý návrh vychází z úspěšné identifikace atributů, které jsou zjistitelné z proudového spektra motoru pohánějícího čerpadlo. Tyto atributy jsou následně využity při konstrukci robustního klasifikátoru transformujícího průběžně vzorkované hodnoty chybových atributů na chybové klasifikace. Řešení této úlohy lze nalézt v [8].

### 4 Závěr

Článek diskutuje teoretické a praktické otázky průmyslové aplikace umělé inteligence se zaměřením na strojové učení. Právě schopnost učit se totiž zvyšuje inteligentní systémy nad systémy ostatní. Úvodní část se snaží stručně naznačit současné možnosti strojového učení, jeho silné stránky a také slabiny. Pokud tyto silné i slabé stránky shrneme do jedné populární věty, lze říci, že na těchto systémech je přitažlivá právě schopnost učení či adaptace a slabinou je, že tyto schopnosti zatím ve většině důležitých ohledů nedosahují schopností lidských. V mnoha směrech je tak návrh nového řešení stále do značné míry kooperativním úsilím znalostního inženýra a experta v dané problémové oblasti. Vývoj v této vědní oblasti, který je v úvodní části také naznačen, směřuje k postupné formalizaci a automatizaci vývoje aplikace a současně k rozšiřování adaptivních schopností systémů samotných.

V následující praktické části je prezentováno konkrétní řešení průmyslové úlohy. Důraz je přitom kladen na globální pohled na návrh celé aplikace a roli strojového učení. Systém SPS pak představuje ukázkou systému integrujícího techniky strojového učení a dobývání znalostí z dat.

### Literatura

1. Atkeson, C., G., Moore, A., W. and Schaal, S.: Locally weighted learning. Artificial Intelligence Review, special issue on Lazy Learning, 1997.
2. Anguita, D.: Smart Adaptive Systems: State of the Art and Future Directions of Research. In: Final Programme & Proceedings. Aachen : Verlag Mainz, 2001, vol. 1, p. 18.
3. Brazdil, P., Soares, C.: A Comparison of Ranking Methods for Classification Algorithm Selection. In Proceedings of the 11<sup>th</sup> European Conference on Machine Learning, p. 63-74, Springer, 2000.
4. Bůcha, J.: Strojové učení. Vesmír 9/2001. p. 503–505.
5. Frank, E., Hall, M.: A Simple Approach to Ordinal Classification. In Machine Learning: proceedings ECML 2001, De Raedt, Flach (eds.), Springer Verlag, 2001.
6. Gunopulos, D., Das, G.: Time Series Similarity Measures. In Tutorial Notes of the Association for Computing Machinery Sixth International Conference on Knowledge Discovery and Data Mining, 2000, p. 243-307.
7. Kittler, J., Hatef, M., Duin, R., Matas, J.: On Combining Classifiers. IEEE Transactions on Pattern Analysis and Machine Intelligence, vol. 20/3, p. 226 – 239, 1998.
8. Kléma, J.: Prototype Applications of Instance-Based Reasoning. PhD Thesis, Czech Technical University, Department of Cybernetics, 2002.
9. Kléma, J., Kout, J.: Prediction of Gas Consumption. In: Systems Integration '99. Prague : University of Economics, 1999, p. 119-128. 1999.
10. Kohavi, R., Provost, F.: Glossary of Terms. Special Issue on Applications of Machine Learning and the Knowledge Discovery Process. Machine Learning, 30, p. 271-274, Kluwer Academic Publishers, Boston, 1998.
11. Kolodner, J.: Case-Based Reasoning, Morgan Kaufmann Publishers, 1996.
12. Mitchell, T.: Machine Learning. Morgan Kaufmann, 1997.



## **Poděkování**

Tento článek vznikl za podpory výzkumného záměru Rozhodování a řízení pro průmyslovou výrobu, MSM 212300013. Projekt je zaměřen na komplexní rozvoj prostředků pro podporu inteligentního rozhodování a řízení na všech úrovních podniků a na moderní metody a techniky automatického řízení výrobních procesů.