

Uchazeč: doc. Ing. Petr Červa, Ph.D.

Podpis:

Hodnocené období: 2021–2025

Poznámka: Tabulky lze přidáním řádků podle potřeby upravit. Doporučujeme uvádět maximálně pět výsledků daného typu.

A1. Vědecko výzkumná činnost

Základní výzkum (hodnocený především na základě publikací nových poznatků)
1. výsledek Matějů, L., Nouza, J., Červa, P., Žďánský, J.: Combining multilingual resources to enhance end-to-end speech recognition systems for Scandinavian languages. Speech Communication, 170, art. 103221, 2025.
Charakterizace Článek řeší dlouhodobě otevřený problém tvorby moderních end-to-end systémů rozpoznávání řeči pro jazyky s omezenými trénovacími zdroji, konkrétně pro dánštinu, švédštinu a norštinu. Hlavním přínosem je systematické prozkoumání způsobů, jimiž lze pro tyto jazyky využít existující vícejazyčné zdroje, ať už v podobě řečových dat příbuzných jazyků, nebo již natrénovaných modelů. Z porovnávání variant se jako nejúčinnější ukázala inicializace parametrů enkodéru z jiného dostupného modelu, označovaného v článku jako dárcovský. Zásadním zjištěním je, že tento postup přináší zisk nejen při malém objemu dat cílového jazyka, ale i tehdy, jsou-li k dispozici tisíce hodin nahrávek, a dokonce i v případě, kdy dárcovský model pochází ze vzdáleného jazyka. Práce podrobně analyzuje vliv volby dárcovského jazyka, velikosti trénovacích dat na straně cílového i dárcovského modelu a možnost souběžného využití více dárcovských modelů. Tento poznatek byl dále rozvinut do praktického postupu sběru dat, v němž několik dárcovských modelů běží paralelně a vzájemně slouží jako kontrolní mechanismus. Tím se daří odhalovat anotační chyby v trénovacích sadách i při automatizovaném vytěžování dat z veřejných zdrojů, jako jsou televizní vysílání, parlamentní záznamy, YouTube, podcasty nebo audioknihy. Součástí práce je rovněž tvorba rozsáhlých testovacích sad, které byly zveřejněny a umožňují reprodukovatelné srovnání. Výsledný systém byl porovnán s komerčními službami předních poskytovatelů rozpoznávání řeči pro skandinávské jazyky a dosahuje v průměru vyšší přesnosti. Článek tak přináší nejen metodický poznatek o přenositelnosti akustických reprezentací mezi jazyky, ale i ověřený postup použitelný pro další jazyky s omezenými zdroji.
2. výsledek Poláček, M., Červa, P., Žďánský, J.: Lightweight online punctuation and capitalization restoration for streaming ASR systems. Speech Communication, 173, art. 103269, 2025.
Charakterizace Práce se zabývá automatickým doplňováním interpunkce a velkých písmen do výstupu systémů rozpoznávání řeči, tedy úlohou, bez jejíhož vyřešení zůstává automatický přepis obtížně čitelný. Novost přístupu spočívá v tom, že je navržen pro provoz v reálném čase, s velmi krátkou dobou odezvy a latencí pouhých čtyř slov, což jej činí použitelným pro živé titulkování televizního a rozhlasového vysílání. Metoda pracuje pouze s textovým vstupem, čímž se odlišuje od konkurenčních řešení kombinujících textové a akustické příznaky, a přesto je s nimi v článku přímo porovnávána na prepisech rozhlasových diskusí a televizních debat. Výsledné řešení je postaveno na dvojici navazujících modelů architektury ELECTRA-small doplněných jednoduchými klasifikačními hlavami, kde první model obnovuje interpunkci a druhý provádí kapitalizaci. Tato dvoustupňová dekompozice úlohy je klíčová z hlediska výpočetní nenáročnosti, protože umožňuje použít malé modely namísto jednoho velkého. Systém obnovuje otazníky, čárky, tečky a velká písmena a byl vyhodnocen pro češtinu i němčinu, čímž je doložena jeho jazyková přenositelnost. Součástí srovnání je i existující systém pro angličtinu, což umožňuje zasadit dosažené výsledky do mezinárodního kontextu. Autoři zároveň zveřejnili data použitá pro vyhodnocení a testování, což přispívá k reprodukovatelnosti a dalšímu rozvoji této úlohy. Prvním autorem článku je Ing. Martin Poláček, současný doktorand uchazeče, což dokládá zapojení doktorandů do publikování v impaktovaných časopisech. Přínos práce je tedy dvojitý: metodický, v podobě architektury umožňující vysokou kvalitu při nízké latenci, a praktický, neboť výsledek je přímo nasaditelný v provozních systémech pro přepis řeči.

3. výsledek

Kynych, F., Červa, P., Žďánský, J., Svendsen, T. K., Salvi, G.: A lightweight approach to real-time speaker diarization: from audio toward audio-visual data streams. EURASIP Journal on Audio, Speech, and Music Processing, 2024, art. 62.

Charakterizace

Článek se věnuje diarizaci mluvčích, tedy úloze určit, kdo a kdy hovoří, v režimu zpracování živého datového proudu. Online diarizace je podstatně obtížnější než offline varianta, protože systém nemá k dispozici celý záznam a musí rozhodovat průběžně, při zachování nízké latence a omezených výpočetních nároků. Navržený přístup je záměrně koncipován jako výpočetně nenáročný, což umožňuje jeho nasazení v reálných systémech pro přepis vysílání bez potřeby rozsáhlé výpočetní infrastruktury. Podstatným rozšířením oproti dosavadním pracím je postup od čistě audio zpracování směrem k audiovizuálním datovým proudům, kdy jsou vedle akustické informace využívány i obrazové podněty. Kombinace obou modalit umožňuje zpřesnit určení mluvčího zejména v situacích, kdy je akustický signál zarušen nebo kdy se promluvy překrývají. Práce vznikla v mezinárodní spolupráci s Norwegian University of Science and Technology v Trondheimu a je výstupem projektu NordTrans. Metodicky navazuje na dřívější práce autorského týmu v oblasti detekce změny mluvčího a adaptace akustických modelů. Prvním autorem je doktorand uchazeče, což dokládá zapojení doktorandů do mezinárodně publikovaného výzkumu. Výsledky jsou přímo využitelné v úloze automatického přepisu, kde diarizace slouží jako vstup pro adaptaci rozpoznávacího systému na konkrétního mluvčího.

4. výsledek

Červa, P., Matějů, L., Žďánský, J., Šafařík, R., Nouza, J.: Identification of related languages from spoken data: moving from off-line to on-line scenario. Computer Speech and Language, 68, art. 101180, 2021.

Charakterizace

Hlavním přínosem článku je návrh nové metody pro online identifikaci jazyka, která ve výsledku umožňuje přepis datových proudů obsahujících promluvy ve více jazycích. Jde o téma, které bylo do té doby jen málo probádané, zejména pro dvojice vzájemně blízkých jazyků, u nichž je rozlišení podstatně obtížnější než u jazyků nepříbuzných. Jednotlivé aspekty metody jsou v článku doloženy řadou navazujících experimentů, počínaje offline identifikací pro jedenáct hlavních slovanských jazyků, přes online identifikaci na uměle připraveném vícejazyčném datasetu, až po online detekci jazyka ve skutečných televizních a rozhlasových pořadech obsahujících směs češtiny a slovenštiny. Metoda zpracovává data po jednotlivých segmentech, kdy vstupem je vícejazyčný proud řečových segmentů a výstupem posloupnost odpovídajících jazykových značek. Jako dekodér je využíván vážený převodník s konečným počtem stavů, který vyhlazuje výstup jazykového klasifikátoru a potlačuje tak náhodné přepnutí jazyka. Klasifikátor pracuje nad vícejazykovými bottleneck příznaky, přičemž jak tyto příznaky, tak samotný klasifikátor jsou založeny na architekturách hlubokých neuronových sítí umožňujících modelovat dlouhodobější časové závislosti. Získané výsledky ukazují, že navržené schéma určí jazyk promluvy v reálných dvojjazyčných pořadech s průměrnou latencí kolem 2,5 sekundy. Chybovost výsledného přepisu je přitom jen přibližně o tři procenta vyšší než v situaci, kdy by přepínání jazyků vycházelo z referenčních, člověkem vytvořených jazykových značek. Práce tak prokazuje, že plně automatické zpracování vícejazyčného vysílání je prakticky dosažitelné.

Aplikovaný výzkum (hodnocený především na základě realizací nových technologií, konstrukcí, apod.)	
1. výsledek (realizace)	JUSEARCH – softwarový modul pro prohledávání justičních rozhodnutí přirozeným jazykem (2025, podíl TUL 100 %). Autoři: Petr Červa, Lukáš Matějů, Martin Poláček.
Charakterizace (V-V přínos, uplatnění, patent, osobní podíl, ...)	JUSEARCH je sémantické vyhledávací jádro nad korpusem soudních rozhodnutí, které umožňuje položit dotaz v přirozeném jazyce, a to i bez diakritiky nebo s překlepy, a nalézt věcně relevantní dokumenty i tam, kde se dotaz a text rozhodnutí lexikálně nepřekrývají. Modul provádí souběžné vyhledávání dvěma vzájemně komplementárními vektorizačními modely, jejichž výstupy jsou kombinovány metodou Reciprocal Rank Fusion a doplňovány o vyhledávání podle spisových značek. Originalita řešení spočívá v tom, že oba modely jsou metodou LoRA doladěny na doménovém korpusu dvojic otázka–úsek judikatury. Takový korpus pro českou judikaturu neexistoval a bylo nutné jej vytvořit; v případě Ústavního soudu ČR vznikl během pilotního provozu z 2 000 dvojic, jejichž věcnou správnost posoudili právníci soudu, nikoli autoři systému. Právě tento krok je metodicky podstatný, neboť tazatel formuluje dotaz jinými slovy, než jakými je psána judikatura, a model naučený pouze na textech rozhodnutí se tuto mezeru nenaučí překlenout. Přínos byl kvantifikován na testovací sadě 300 reálných uživatelských dotazů: produkční konfigurace dosahuje nDCG@10 na úrovni 0,53 oproti 0,44 u obecného multilingválního modelu, tedy o 21 % relativně lépe, a podíl případů, kdy hledaný úsek nepronikl do předávaného kontextu, klesl z deseti na dvě procenta. Klasické lexikální vyhledávání metodou BM25 přitom zaostává za všemi vektorovými variantami. Index Ústavního soudu ČR obsahuje 104 870 dokumentů rozdělených do 766 378 úseků, v nichž modul díky algoritmu HNSW nalezne kontext do 200 ms. Výsledek je prostřednictvím spin-off společnosti nanotrix.ai nasazen v ostrém provozu na Ústavním soudu České republiky od 23. února 2026, kde podle vyjádření generálního sekretáře soudu zodpověděl již více než 36 000 dotazů, pracuje spolehlivě, bez výpadků a s nízkou mírou halucinací; od 1. července 2026 je nasazen rovněž na Ústavním soude Slovenskej republiky nad více než 50 000 rozhodnutími. Podle vyjádření soudu rozšířilo nasazení okruh osob schopných se v judikatuře prakticky orientovat a zrychlilo vyhledávání i u profesionálních uživatelů; podle vedoucí analytického týmu Ústavního soudu ČR jde o jedno z prvních nasazení tohoto typu v Evropě. V současné době probíhá jednání o nasazení nad agendou veřejného ochránce práv ČR. Osobní podíl uchazeče spočívá ve vedení vývoje a v návrhu architektury vyhledávací vrstvy; spoluautory výsledku jsou Ing. Lukáš Matějů, Ph.D., jeho bývalý doktorand, a Ing. Martin Poláček, jeho současný doktorand.
2. výsledek (projekt)	FY01010017 LibriBot – Doporučování a vyhledávání literárních děl pomocí umělé inteligence (2025–2027, program TWIST, Ministerstvo průmyslu a obchodu; hlavní příjemce Tritius Solutions, a.s.)
Charakterizace (V-V přínos, uplatnění, patent, osobní podíl, ...)	Cílem projektu je vývoj dvou softwarových služeb využívajících umělou inteligenci pro doporučování a vyhledávání literárních děl v prostředí knihoven a archivů. Služba pro obsahové doporučování staví na vektorové reprezentaci díla vypočtené z jeho metadat pomocí vlastního jazykového modelu, druhá služba umožňuje volné dotazování na obsah fondu metodou Retrieval-Augmented Generation. Obě služby podporují češtinu i slovenštinu. Výzkumná nejistota spočívá zejména v tom, že pro české a slovenské prostředí dosud neexistuje doporučovací řešení opírající se o vektorovou reprezentaci obsahu, a je tedy nutné vytvořit dostatečně kvalitní vektorizační modely i ověřit optimální způsob využití metadat a vyhledávací strategie nad katalogy čítajícími stovky tisíc svazků. Uchazeč je řešitelem za TUL a vede liberecký tým, který v projektu odpovídá za výzkumné jádro, tedy za tvorbu a optimalizaci vektorizačních modelů a za vývoj obou vyhledávacích a doporučovacích modulů; partner zajišťuje databázi metadat a integraci do své platformy. Na TUL připadá 6 860 tis. Kč z celkových uznaných nákladů 14 106 tis. Kč. Zájem o výsledné služby potvrdily velké knihovny v České republice i na Slovensku; po dokončení mohou být dostupné stovkám knihoven na obou trzích.
3. výsledek (realizace)	Softwarové moduly pro automatický přepis a zpracování mluvené norštiny (2024, výstup projektu TO01000027 NordTrans, podíl TUL 100 %). Autoři: Petr Červa, Jan Nouza, Jindřich Žďánský, Lukáš Matějů.
Charakterizace (V-V přínos, uplatnění, patent, osobní podíl, ...)	Výsledkem je modulární softwarová technologie převádějící mluvenou norštinu do textové podoby. Jádro tvoří hluboká neuronová síť typu transformer natrénovaná originálně vyvinutým postupem s využitím vícejazyčných dat, jejíž výstupní hypotézy jsou v reálném čase vyhodnocovány dekodérem; navazuje postprocessingový modul a

modul doplňující interpunkci a velká písmena pomocí předtrénovaného jazykového modelu. Specifickou vlastností řešení je podpora obou psaných standardů norštiny, bokmål i nynorsk, čímž technologie přispívá k uchování jazykové rozmanitosti a je použitelná například pro přepis parlamentních jednání, kde se oba standardy běžně vyskytují. Technologie je výsledkem několikaletého výzkumu v rámci mezinárodního projektu NordTrans, jehož dílčí poznatky byly publikovány ve čtyřech odborných pracích, z toho ve dvou článcích v časopise Speech Communication. Srovnání provedené na veřejně dostupných datových sadách zahrnujících nahrávky z parlamentu, televize, platformy YouTube a dalších zdrojů prokázalo, že vyvinutá technologie dosahuje pro norštinu v průměru nejvyšší přesnosti ze všech porovnávaných řešení, včetně open-source modelu Whisper společnosti OpenAI a komerčních cloudových služeb společností Google, Microsoft a Speechmatics. Software umožňuje online i offline přepis s latencí několika sekund a paměťovými nároky v řádu jednotek gigabajtů, je tedy provozovatelný bez rozsáhlé výpočetní infrastruktury. Výsledek je komerčně využíván partnerem projektu, společností NEWTON Technologies, a.s., a to v platformě Beey pro offline přepis vysílání a jednání a pro tvorbu titulků pro zákazníky ze skandinávské oblasti, a v platformě BeeyLive pro živé titulkování konferencí a akcí konaných v norštině. Využití technologie v režimu živého titulkování je doloženo dopisem Norského velvyslanectví adresovaným společnosti NEWTON Technologies. Vedle komerčního uplatnění má výsledek i společenský rozměr: umožňuje inkluzivní přístup k obsahu pro osoby se sluchovým postižením prostřednictvím automatického titulkování, podporuje digitalizaci záznamů veřejné správy a je využitelný pro výuku a analýzu mluveného jazyka. Podíl TUL na výsledku je stoprocentní. Osobní podíl uchazeče spočívá ve vedení řešitelského týmu TUL a v koordinaci výzkumných prací napříč konsorciem, jehož součástí byla vedle komerčního partnera i Norwegian University of Science and Technology v Trondheimu.

4. výsledek (realizace)

Multilingvální softwarová technologie pro detekci a včasné upozornění (2021, výstup projektu TH03010018 DeepSpot, podíl TUL 50 %). Autoři: Petr Červa, Jan Nouza, Jan Václ, Lenka Weingartová.

Charakterizace (V-V přínos, uplatnění, patent, osobní podíl, ...)

Platforma využívá hluboké neuronové sítě a umožňuje automaticky detekovat klíčová slova v mluvené řeči a přepisovat plynulou řeč ve třinácti hlavních slovanských jazycích, konkrétně v češtině, slovenštině, polštině, ruštině, ukrajinštině, běloruštině, slovinštině, srbštině, chorvatštině, bosenštině, černohorštině, makedonštině a bulharštině. Vedle zpracování řeči zajišťuje rovněž detekci entit v obrazových datech, například komerčních log nebo osob podle obličeje, a je složena z jazykově závislých i nezávislých modulů zpracovávajících multimediální data na signálové, modelové, detekční a prezentační úrovni. Pracuje v online režimu nad živým proudem dat i v offline režimu s velmi rychlou detekcí v datovém archivu. Výsledek byl v hodnocení podle Metodiky M17+ ohodnocen stupněm 2, tedy jako výsledek na vynikající úrovni s významným dopadem na společnost nebo trh. Platforma byla hlavním výstupem čtyřletého projektu TA ČR TH03010018 s celkovými uznanými náklady 29 765 tis. Kč a jednotlivé moduly vznikly na základě původního výzkumu prezentovaného v sedmnácti příspěvcích na mezinárodních konferencích indexovaných ve WoS nebo Scopus a v jednom článku v impaktovaném časopise kvartilu Q2. Podíl TUL na výsledku je padesátiprocentní, přičemž na TUL probíhaly hlavní výzkumné a vývojové práce, doložené autorstvím uvedených publikací; partner se soustředil na sběr dat, testování a integraci modulů do výsledných služeb. Na základě platformy provozuje společnost NEWTON Technologies řadu nadstavbových služeb, zejména webovou platformu Beey pro automatický přepis a tvorbu titulků, variantu Beey workflow pro korektury automatických přepisů a službu TVR Alerts upozorňující klienty na výskyt zadaných klíčových slov v televizním a rozhlasovém vysílání. Mezi uživatele patří Česká národní banka, TV NOVA, Czech News Center, DVTV a řada evropských monitoringových agentur v Polsku, na Slovensku, v Chorvatsku, Srbsku, Bosně a Hercegovině, Severní Makedonii, Černé Hoře a na Ukrajině; využívání technologie je doloženo osvědčeními vybraných klientů. Přesnost automatického přepisu je plně srovnatelná se službami dostupnými v cloudových portálech velkých nadnárodních společností, což dokládá i to, že zákazníci partnera dávají přednost tomuto řešení před nástroji postavenými na těchto cloudových službách. Platforma byla v navazujícím mezinárodním projektu NordTrans rozšířena o podporu norštiny, švédštiny a dánštiny. Osobní podíl uchazeče spočívá ve vedení řešitelského týmu TUL a v odpovědnosti za moduly akustického a jazykového modelování.

A2. Pedagogická a vzdělávací činnost

Přednášková činnost (garance a vedení přednášek)
Bakalářský studijní program Informační technologie, specializace Inteligentní systémy: Úvod do strojového učení — tvůrce předmětu, přednášející a cvičící. Metody vytěžování dat — garant.
Navazující magisterský studijní program Informační technologie: Úvod do počítačové lingvistiky — zakladatel předmětu, garant a přednášející. Pokročilé programování na platformě Java — spoluzakladatel předmětu (2016) a spoluvyučující; výuka ukončena v roce 2025. Vybrané statě ze strojového učení — nový předmět zavedený v rámci akreditace programu v roce 2025; uchazeč jej přednáší a vede cvičení. Metody zpracování přirozeného jazyka — nový předmět zavedený v rámci akreditace programu; uchazeč jej přednáší a vede cvičení. Cloudové technologie I, Cloudové technologie II — garant.
Doktorský studijní program Technická kybernetika: Hluboké neuronové sítě — spolugarant předmětu; o předmět projevují zájem i studenti z jiných fakult TUL.
Učebnice a výukové pomůcky (charakteristika učebnice, výukové pomůcky)
Individuální vzdělávací činnost (vedení projektu, diplomové práce, doktoranda, kvantitativní i kvalitativní hodnocení)
Doktorandi Dokončené PhD: Ing. František Kynych, Ph.D. (Fakulta mechatroniky, informatiky a mezioborových studií TUL; školitel; obhájeno 2025). Probíhající PhD: Ing. Martin Poláček (Fakulta mechatroniky, informatiky a mezioborových studií TUL; školitel; po obhajobě tezí). Probíhající PhD: Ing. Martin Motejlek (Fakulta mechatroniky, informatiky a mezioborových studií TUL; školitel; 2. ročník studia). Mimo hodnocené období: Ing. Lukáš Matějů, Ph.D. (školitel; obhájeno 2020). Doktorandi a absolventi uchazeče jsou spoluautory jeho časopiseckých publikací i aplikovaných výsledků laboratoře. Ing. Kynych je prvním autorem článku o diarizaci mluvcích publikovaného v EURASIP Journal on Audio, Speech, and Music Processing a absolvoval výzkumnou stáž ve skupině prof. G. Salviho na NTNU v Trondheimu; Ing. Poláček je prvním autorem článku v časopise Speech Communication a spoluautorem výsledku JUSEARCH; Ing. Matějů je prvním autorem dalšího článku v témže časopise a rovněž spoluautorem výsledku JUSEARCH.
Studentské projekty 2021 semestrální projekt (bakalářský studijní program): Mariia Buntovskikh — Klasifikace obrázků Lego minifigurek. 2021 semestrální projekt (navazující magisterský studijní program): Bc. Vratislav Dutka — Webová aplikace pro on-line tvorbu řečových databází.

Bakalářské a diplomové práce 2022 BP: David Šafařík — Tvorba modelů pro přepis řeči v italštině. 2023 DP: Bc. Martin Poláček — Automatické generování interpunkce v systémech rozpoznávání řeči (oceněna Cenou rektora TUL). Téma práce uchazeč dále rozvinul do společných publikací na konferenci INTERSPEECH a v časopise Speech Communication; autor po obhajobě pokračuje v doktorském studiu pod vedením uchazeče. 2025 DP: Bc. Martin Motejlek — Data Augmentation Methods for Improving Performance of E2E ASR Systems on Distant Speech (práce obhájena v anglickém jazyce, oceněna Cenou děkana FM). Autor po obhajobě pokračuje v doktorském studiu pod vedením uchazeče.
Oponentská činnost 2021: oponent disertační práce a člen komise pro její obhajobu — Ing. Marie Kunešová „Diarizace řečníků“, Fakulta aplikovaných věd, Západočeská univerzita v Plzni. 2026 (mimo hodnocené období): oponent habilitační práce Ing. Ladislava Lence, Ph.D. „Analysis of Text and Image Documents“, Fakulta aplikovaných věd, Západočeská univerzita v Plzni.

Podíl na garantování Bc., Mgr. a Ph.D. oboru (přínos k profilu absolventa)
Bakalářský studijní program Informační technologie — specializace Inteligentní systémy Uchazeč je autorem myšlenky a koncepce specializace Inteligentní systémy. Specializace je vystavěna ze dvou vzájemně navazujících linií předmětů. První pokrývá teorii zpracování signálů, dat a strojového učení v posloupnosti Signály a informace, Strojové učení, Aplikace neuronových sítí, Úvod do zpracování obrazu, Metody vytěžování dat a Modelování signálů a dat. Druhá linie je zaměřena na softwarové technologie pro zpracování velkých dat a podporu strojového učení v posloupnosti Databáze pro BigData a Technologie pro BigData. Absolvent specializace tak získává jak metodické zázemí, tak schopnost navržená řešení realizovat a provozovat.
Navazující magisterský studijní program Informační technologie Od roku 2025 je uchazeč garantem navazujícího magisterského studijního programu Informační technologie, který v nové podobě připravil a akreditoval. Nešlo o formální obnovení akreditace, nýbrž o přepracování celého obsahu programu. Po vzoru bakalářského stupně byly zavedeny upravené specializace Inteligentní systémy, Softwarové inženýrství a Vestavné systémy, takže na sebe oba stupně studia navazují a student může zvolené zaměření dále prohlubovat. Přepracována byla rovněž výuka matematiky, do níž byla nově zařazena algebra. Předmět Vybrané statě ze strojového učení byl zaveden jako povinný pro všechny specializace, a to s ohledem na pronikání metod umělé inteligence do tvorby softwaru: absolvent programu se s nimi setká bez ohledu na zvolené zaměření, a musí jim proto rozumět.

A3. Ostatní významné aktivity

Výkon funkce
Vedoucí Laboratoře umělé inteligence (AILab@TUL), od roku 2021
Přínos pro rozvoj vedeného pracoviště (týmu) Uchazeč vede výzkumný tým o osmi členech, z toho dva doktorandy. V roce 2024 bylo zaměření pracoviště rozšířeno ze zpracování mluvené řeči na umělou inteligenci obecněji, což se promítlo do nového názvu laboratoře. Uchazeč dlouhodobě zajišťuje pro laboratoř projektové financování. V hodnoceném období šlo o tyto projekty, uvedené podle uznaných nákladů: TH03010018 DeepSpot (2018–2021, hlavní příjemce NEWTON Technologies, a.s.): celkové uznané náklady 29 765 tis. Kč, z toho na TUL 15 264 tis. Kč, tedy 51 %. TO01000027 NordTrans (2021–2024, mezinárodní projekt s NTNU Trondheim): uznané náklady připadající na TUL 12 387 tis. Kč, hrazené z 9 228 tis. Kč z národních zdrojů a 3 159 tis. Kč z veřejných zahraničních zdrojů. FY01010017 LibriBot (2025–2027, hlavní příjemce Tritius Solutions, a.s.): celkové uznané náklady 14 106 tis. Kč, z toho na TUL 6 860 tis. Kč, tedy 49 %. Uchazeč je řešitelem za TUL a vede liberecký tým odpovědný za výzkumné jádro projektu. FY01010005 Pokročilá optimalizace vývoje ozubených kol s využitím metod umělé inteligence (2025–2027, hlavní

příjemce LENAM, s.r.o.): celkové uznané náklady 9 188 tis. Kč, z toho na TUL 4 461 tis. Kč, tedy 49 %. Uchazeč se zásadním způsobem podílel na přípravě projektu a je členem řešitelského týmu; celkový úvazek jeho laboratoře na projektu činí 1,1. Z projektů, v nichž byl řešitelem za TUL, zajistil uchazeč pro pracoviště prostředky ve výši přibližně 34,5 mil. Kč. Programy, do nichž byly projekty podány, vyžadují jako hlavního uchazeče podnik, takže TUL v nich vystupuje formálně jako další účastník. Případá na ni však ve všech případech zhruba polovina uznaných nákladů a věcně odpovídá za výzkumné pracovní balíčky, zatímco partner za integraci a uvedení na trh. Uvedené projekty zároveň rozšířily aplikační záběr pracoviště směrem k vyhledávacím a doporučovací systémům, ke konverzačním agentům nad rozsáhlými dokumentovými databázemi a nově i do oblasti strojírenského návrhu. Dalším kanálem přenosu výsledků a spolupráce s praxí je univerzitní spin-off nanotrix.ai.

Výkon funkce
Jednatel spin-off společnosti nanotrix.ai, od roku 2025
Přínos pro rozvoj vedeného pracoviště (týmu)
Prostřednictvím univerzitního spin-offu je zajištěn přenos výsledků laboratoře do praxe formou licencování technologií vyvinutých na TUL. Tímto kanálem bylo realizováno nasazení modulu JUSEARCH na Ústavním soudu České republiky a Ústavném soude Slovenskej republiky. Spin-off zároveň vytváří uplatnění pro absolventy doktorského studia a udržuje zpětnou vazbu mezi potřebami praxe a výzkumným programem laboratoře. Vedle výsledků licencovaných od TUL vyvíjí a provozuje společnost i vlastní produkty stavějící na témže metodickém základu: konverzační a vyhledávací asistenti jsou nasazení na webech zhruba padesáti měst a městských částí, v jejichž správních obvodech žije přibližně 1,4 milionu obyvatel, a dále na webech velkých nemocnic. Měřeno počtem obyvatel obcí, jejichž weby jsou tímto způsobem pokryty, jde v segmentu české veřejné správy o nejrozšířenější řešení svého druhu. Tato tržní pozice je pro laboratoř podstatná ze dvou důvodů: potvrzuje, že výzkumné směry, jimž se pracoviště věnuje, odpovídají reálné poptávce, a zároveň poskytuje přístup k provozním datům a k uživatelské zpětné vazbě v měřítku, jakého by samotné akademické pracoviště nedosáhlo. Touto cestou je komercializován také VoiceBot, tedy konverzační agent spojující sémantické vyhledávání s rozpoznáváním řeči, který umožňuje vyhledávat i prostřednictvím telefonní linky. Tento výsledek vznikl v roce 2026, a proto není uveden mezi aplikovanými výsledky za hodnocené období; z hlediska odborného zaměření pracoviště je však podstatný tím, že propojuje hlavní dřívější téma pracoviště, tedy zpracování mluvené řeči, s hlavním současným tématem, jímž je vyhledávání a konverzační přístup k rozsáhlým dokumentovým databázím. VoiceBot je od roku 2026 nasazen v ostrém provozu v Děčíně a na Praze 11; připravuje se jeho nasazení na Praze 2, 13 a 14, v Hradci Králové a v Jihlavě.

Členství (ve vědeckých radách, v radách redakčních časopisů, funkce ve vědeckých společnostech atd.)
Člen Akademického senátu Fakulty mechatroniky, informatiky a mezioborových studií TUL od roku 2023, druhé volební období.
Osobní přínos
Podíl na samosprávě fakulty, zejména na projednávání vnitřních předpisů, rozpočtu a záměrů v oblasti studijních programů. Opakované zvolení dokládá důvěru akademické obce fakulty.

Jiné aktivity
Recenzní činnost
Přínos
Pravidelný recenzent příspěvků pro konferenci Interspeech, která je hlavním mezinárodním fórem v oboru zpracování mluvené řeči; příležitostně rovněž recenzent pro odborné časopisy.

Jiné aktivity
Hodnotitelská činnost
Přínos
Hodnotitel projektových návrhů pro Technologickou agenturu ČR v programu TREND (2023, pět posudků) a pro agenturu CzechInvest (2026, mimo hodnocené období).
Jiné aktivity
Mezinárodní spolupráce
Přínos
Dlouhodobá spolupráce s Norwegian University of Science and Technology v Trondheimu navázaná v rámci projektu NordTrans, jejímž výsledkem jsou společné publikace a výzkumná stáž doktoranda uchazeče ve skupině prof. Giampiera Salviho. Opakované výukové pobyty na University of Granada v letech 2009 a 2026 v rámci programu Erasmus. Na tomtéž pracovišti absolvovali pobyt rovněž doktorandi uchazeče, Ing. Lukáš Matějů, Ph.D. a Ing. Martin Poláček, takže nejde o kontakt vázaný na jednu osobu, ale o spolupráci využívanou celým týmem.