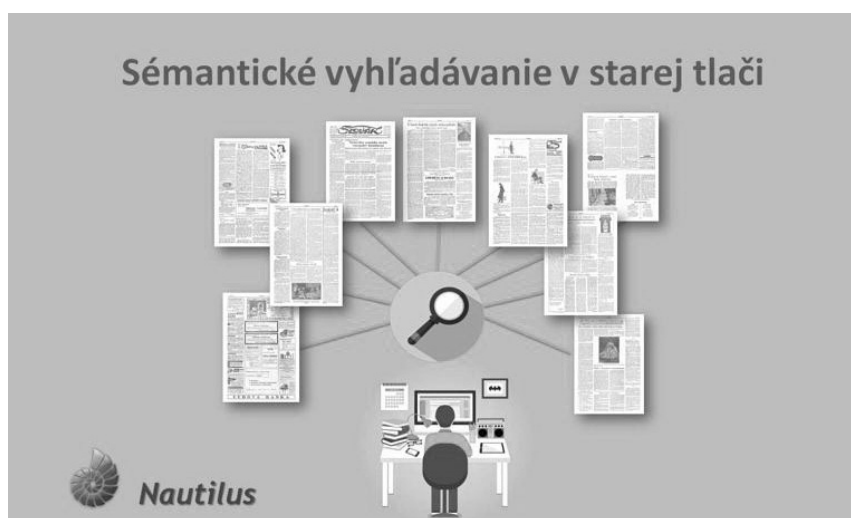


NAUTILUS – Pomáhame sprístupniť informácie už raz nájdené



Nemusíte byť študentom, aby ste potrebovali nájsť nejaké informácie. My v škole pracujeme prevažne s digitálnymi knižnicami. S radosťou sme preto prijali správu o tom, že vďaka digitalizácii máme dnes len v Univerzitnej knižnici k dispozícii viac ako milión strán z periodík. Lenže o to väčšie bolo naše sklamanie, keď sme zistili, že vyhľadávanie v plných textoch je síce fajn, ale tie desiatky strán, ktoré získame, nie sú žiadnou výhrou. Chcelo by to analytický rozpis článkov. Lenže je to vôbec možné? Ved' na jednej takej strane môže byť hneď niekoľko článkov. Tieto a podobné otázky sme si na začiatku kládli aj my. Rozhodli sme sa však využiť naše znalosti a skúsenosti s technológiami a výsledkom je



Obr. 1

Dôvody a východiská projektu

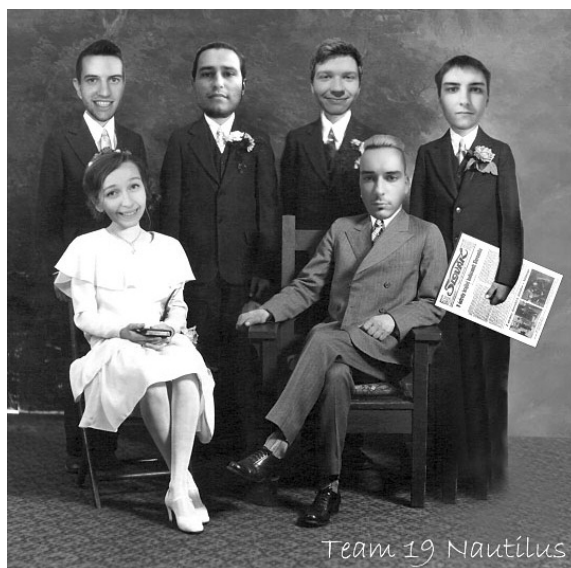
V dôsledku webu 2.0, ale aj rozvoja ďalších informačno-komunikačných technológií sa neustále zrýchľuje nárast dát, a to nielen v textovej podobe. Na druhej strane knižnice a iné inštitúcie tohto typu majú vo svojich fondoch obrovské množstvá dát, ktoré sú ešte aj teraz pre všetky technológie prakticky neprístupné (neviditeľné). Jedná sa predovšetkým o staršie periodiká, ktoré sú spracované len na úrovni súborného záznamu periodika. To, o čom sa píše v jednotlivých článkoch je však pre počítače úplne neviditeľné, aj keď v mnohých prípadoch ide o veľmi cenné informácie a poznatky. Vyhľadávať v nich môžeme len manuálne, teda môžeme čítať obsahy jednotlivých čísiel a konkrétne články a len tak môžeme nájsť tie informácie, ktoré práve potrebujeme alebo môžeme tak získať pre nás zatiaľ úplne nové poznatky. V dôsledku masívnej digitalizácie, ktorá sa u nás začala realizovať pred pár rokmi sa podarilo časť týchto zdrojov dostať do digitálnej podoby, takže sa stali aj pre počítače prístupné. Sprístupnenie a vyhľadávanie sú však dve úplne rozličné veci. Aj keď sú tieto dáta v digitálnej podobe, stále chýbajú nástroje, ktoré by umožňovali rýchlo, odkiaľkoľvek a z akéhokoľvek zariadenia nájsť to, čo práve potrebujeme. Oveľa lepšie je to s výberovými periodikami a aj to len z časti, pretože tu niektoré inštitúcie pristúpili k spracovaniu periodík aj na úrovni analytického rozpisu. To znamená, že bibliografický záznam, vytvárajú nielen pre „časopis“ ako taký, ale tiež pre každé číslo periodika a výberovo aj pre určité články z daného čísla. Takže obsah týchto článkov je rovnako prístupný ako knihy, a to bez ohľadu na to, či sa jedná o klasickú formu dokumentu alebo je dokument vydaný v elektronickej podobe.

Preto sme sa rozhodli v našom projekte vyriešiť proces automatizácie spracovania analytického rozpisu tých periodík, ktoré boli zdigitalizované a bol na nich aplikovaný OCR program. Väčšina inštitúcií u nás k tomuto účelu používa sw nástroj známy ako Abby Recognition, ktorý takto spracované dáta ukladá v špeciálnom XML súbore danej verzie. Aplikácia, ktorá je výsledkom nášho projektu, pomáha knihovníkom pri spracovaní analytického rozpisu archivovaných periodík a tým pádom ich sprístupňuje širokému spektru používateľov. Pomáha taktiež vykonávať ďalšie úpravy textu, ktorý je výsledkom OCR tak, aby bolo možné s ním pracovať pri aplikovaní nástrojov umelej inteligencie a umožniť tak aj úplne presné vyhľadávanie na základe významu. Takto spracované periodiká umožnia bežnému používateľovi nielen rýchlejšie nájsť to, čo práve potrebuje, ale následne tiež identifikovať a odhaliť nové poznatky, ktoré tieto staré dokumenty ukrývajú.

Niečo o nás

Tím Nautilus, to sme my, mladí, motivovaní študenti Fakulty informatiky a informačných technológií STU v Bratislave a pod vedením Ing. Nadeždy Andrejčíkovej, PhD., sa snažíme využiť IKT tak, aby priamo pomáhali pri sprístupňovaní kultúrneho

dedičstva. Na tomto projekte sme tiež priamo spolupracovali s jednou z popredných našich inštitúcií v tejto oblasti, a to s Univerzitnou knižnicou v Bratislave, ktorá okrem iného spravuje rozsiahly fond tlačných periodických dokumentov a buduje Súborný katalóg periodík. Univerzitná knižnica v posledných rokoch digitalizovala veľké množstvo týchto starších ročníkov a čísel periodík.

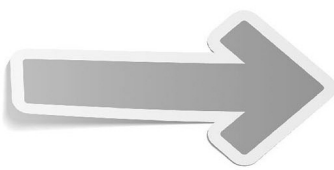


Obr. 2

Spracovanie zdrojov v súčasnosti

Ako sme už uviedli, súčasťou procesu digitalizácie, či skôr tzv. post procesingu, je rozpoznávanie znakov, k čomu sa využíva Adobe Recognition Software. Tento umožňuje zdigitalizované obrazy klasických dokumentov transformovať do špeciálnych XML dokumentov. Takéto súbory, ale aj bežné txt dokumenty, či dokumenty vytvorené v bežných textových editoroch sú tie, na ktoré sa zameriavame a ktoré tvoria základný vstup do nášho projektu.

Tieto špeciálne XML súbory síce obsahujú množstvo informácií, ale sú to zväčša informácie týkajúce sa formátovania a úpravy textu, čo pre účely vyhľadávania nemá dostatočný význam.



```
<?xml version="1.0" encoding="UTF-8"?>
<document version="1.0" producer="FineReader 8.0" xmlns="http://www.abbyy.com/FineReader6-schema-v1.xml" xmlns:xsi="http://www.w3.org/2001/XMLSchema-instance"
xsi:schemaLocation="http://www.abbyy.com/FineReader6-schema-v1.xml http://www.abbyy.com/FineReader6-schema-v1.xml"
pagesCount="28" mainLanguage="Slovak" languages="Slovak,Ci"
page width="3462" height="4986" resolution="400" originalCoords="true">
<block blockType="Text" l="192" t="198" r="586" b="674"><region><rect l="192" t="198"
<text>
<par leftIndent="36">
<line baseline="271" l="244" t="212" r="534" b="270"><formatting lang="Slovak" ff="Tis
<charParams l="244" t="223" r="280" b="270" wordStart="true" wordFromDictionary="false
false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130" charConfidence="31
charParams>
<charParams l="283" t="239" r="312" b="270" wordStart="false" wordFromDictionary="fals
false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130" charConfidence="1
charParams>
<charParams l="316" t="239" r="346" b="270" wordStart="false" wordFromDictionary="fals
false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130" charConfidence="1
charParams>
<charParams l="351" t="238" r="373" b="268" suspicious="true" wordStart="false" wordFr
true" wordNumeric="false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130'
serifProbability="255">s</charParams>
<charParams l="375" t="212" r="398" b="268" suspicious="true" wordStart="false" wordFr
true" wordNumeric="false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130'
serifProbability="255">t</charParams>
<charParams l="406" t="223" r="408" b="225" suspicious="true" wordStart="false" wordFr
true" wordNumeric="false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130'
serifProbability="255">&apos;</charParams>
<charParams l="412" t="228" r="440" b="269" wordStart="false" wordFromDictionary="fals
false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130" charConfidence="1
charParams>
<charParams l="440" t="220" r="473" b="271"> </charParams>
<charParams l="473" t="220" r="510" b="267" wordStart="true" wordFromDictionary="true'
false" wordIdentifier="false" wordPenalty="3" meanStrokeWidth="134" charConfidence="11
charParams>
<charParams l="513" t="236" r="534" b="267" suspicious="true" wordStart="false" wordFr
true" wordNumeric="false" wordIdentifier="false" wordPenalty="3" meanStrokeWidth="134'
serifProbability="100">i</charParams></formatting></line></par>
```

Obr. 3

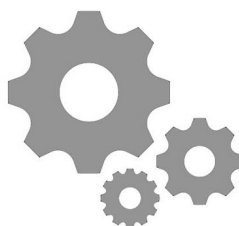
Ako to robíme a náš cieľ

Naším hlavným cieľom v tomto projekte je, ako sme už uviedli, identifikovať tie informácie, ktoré nám umožnia správne rozpoznať jednotlivé bloky textov, ale aj to, ktoré a v akom poradí patria ešte k sebe a tiež to, ktorý z týchto blokov je ten, čo určuje názov, resp. názvovú informáciu. Ako to robíme?

Krok 1: predspracovanie článkov

V prvom kroku, skôr než sme vstupný súbor začali spracovávať, sme museli vykonať viaceré kontroly zamerané na kvalitu vstupného súboru. V tejto časti sme aplikovali viacero metód tak, aby sme mohli spracovateľovi, či administrátorovi poskytnúť čo najpresnejšiu informáciu o nezrovnalostiach, ktoré sme týmito analýzami identifikovali. V prípade, že vstupný súbor spĺňa kvalitatívne kritériá, sme pristúpili k predspracovaniu vstupného súboru. Cieľom tohto predspracovania uvedených vstupných súborov je dosiahnuť vo výslednej forme paragrafy jednotlivých článkov v rovnakých skupinách. K tomuto sme museli vykonať podrobnú analýzu týchto vstupných dát, ktoré sa najčastejšie nachádzajú v rôznych verziách. Identifikovali sme tu viaceré kvalifikátory a iné pravidlá, kombináciou ktorých sa nám podarilo zdefinovať základný algoritmus. Pomocou tohto nášho algoritmu môžeme automatizovať proces extrakcie článkov. Tento algoritmus je založený na rozpoznávaní nadpisov a textov a následnom priradení jednotlivých textov k príslušajúcim nadpisom. Vďaka našej metóde sme už dnes schopní odhaliť a extrahovať až 70 % článkov zo vstupných XML súborov jednotlivých čísel periodík.

```
<?xml version="1.0" encoding="UTF-8"?>
<document version="1.0" producer="FineReader 8.0" xmlns="http://www.abbyy.com/FineReader6-schema-v1.xml" xmlns:xsi="http://www.w3.org/2001/XMLSchema-instance"
xsi:schemaLocation="http://www.abbyy.com/FineReader6-schema-v1.xml http://www.abbyy.com/FineReader6-schema-v1.xml"
pagesCount="28" mainLanguage="Slovak" languages="Slovak,Cz"
page width="3462" height="4986" resolution="400" originalCoords="true">
<block blockType="Text" l="192" t="198" r="586" b="674"><region><rect l="192" t="198"
<text>
<par leftIndent="36">
<line baseline="271" l="244" t="212" r="534" b="270"><formatting lang="Slovak" fff="Tia
<charParams l="244" t="212" r="534" b="270" wordStart="true" wordFromDictionary="false"
false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130" charConfidence="3:
charParams>
<charParams l="283" t="239" r="312" b="270" wordStart="false" wordFromDictionary="false"
false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130" charConfidence="3:
charParams>
<charParams l="316" t="239" r="346" b="270" wordStart="false" wordFromDictionary="false"
false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130" charConfidence="3:
charParams>
<charParams l="351" t="238" r="373" b="268" suspicious="true" wordStart="false" wordFr
true" wordNumeric="false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130"
serifProbability="255"></charParams>
<charParams l="379" t="222" r="398" b="268" suspicious="true" wordStart="false" wordFr
true" wordNumeric="false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130"
serifProbability="255"></charParams>
<charParams l="406" t="223" r="408" b="225" suspicious="true" wordStart="false" wordFr
true" wordNumeric="false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130"
serifProbability="255">&apos;</charParams>
<charParams l="412" t="228" r="448" b="269" wordStart="false" wordFromDictionary="false"
false" wordIdentifier="false" wordPenalty="2" meanStrokeWidth="130" charConfidence="3:
charParams>
<charParams l="440" t="220" r="473" b="271"></charParams>
<charParams l="473" t="220" r="510" b="267" wordStart="true" wordFromDictionary="true"
false" wordIdentifier="false" wordPenalty="3" meanStrokeWidth="134" charConfidence="1:
charParams>
<charParams l="513" t="236" r="534" b="267" suspicious="true" wordStart="false" wordFr
true" wordNumeric="false" wordIdentifier="false" wordPenalty="3" meanStrokeWidth="134"
serifProbability="100"></charParams></formatting></line></par>
```



Obr. 4

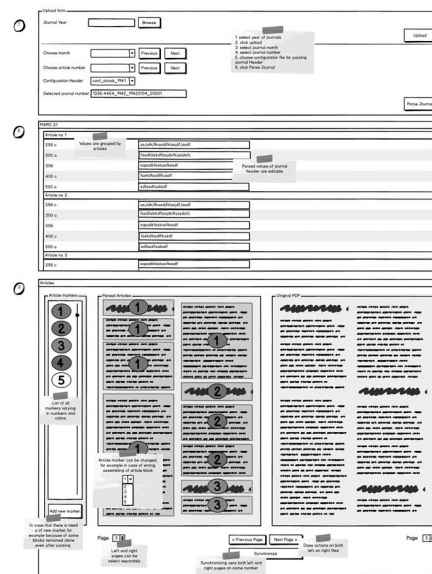
Následne sa snažíme tieto texty spracovať a identifikovať v nich kľúčové slová, neskôr aj význam, v akom boli tieto slová v danom texte použité. Keďže hovoríme o automatizácii procesu spracovania periodík na úroveň ich analytického rozpisu, tak z takto získaných údajov generujeme bibliografické záznamy pre dané číslo periodika, ako aj pre konkrétne články, ktoré sú v ňom obsiahnuté. Samozrejme všetky tieto nami generované bibliografické záznamy článkov obsahujú vždy väzbu na dané číslo periodika a bibliografický záznam čísla periodika, majú v sebe väzbu na existujúci záznam súborného záznamu periodika, prípadne je možné pridať aj väzbu na záznam zviazaného ročníka a zväzku periodika, ale to by si vyžadovalo pri nastavení služby mať k dispozícii tieto bibliografické záznamy zviazaných ročníkov periodík. Bibliografické záznamy generujeme vo formáte MARC21/XML podľa platných katalogizačných pravidiel, aby s nimi mohli priamo pracovať aj knižnično-informačné systémy. K jednotlivým bibliografickým záznamom ukladáme do nášho repozitára aj plné texty a ich vizualizované obrazy tak, aby boli vzájomne prepojené a používateľ ich mohol získať spolu s bibliografickým záznamom.

Krok 2: webová aplikácia

Prídanou hodnotou celého nášho riešenia je priamo webová aplikácia, pomocou ktorej môžu používatelia jednoduchou a intuitívnu formou pracovať s výsledkami nášho riešenia, teda s jednotlivými rozpoznávanými a identifikovanými textami a ich usporiadaním do článkov. Výhodou tejto webovej aplikácie je, že zároveň predstavuje rozhranie pre úpravu samotných plných textov. Keďže kvalita primárnych zdrojov, ako aj schopnosti OCR nástrojov rozpoznávať znaky, nemusia byť a z praxe vieme, že ani nie je vždy stopercentná, je súčasťou nášho projektu tiež návrh a realizácia aplikácie, ktorá umožňuje zamestnancom inštitúcií pre archiváciu takýchto zdrojov, prípadne iným používateľom, editovať aj výsledky OCR spracovania. Umožňuje im teda priamo opravovať daný text, ktorý pri opravách priebežne ukladáme. Rovnako, pomocou tejto aplikácie, umožňujeme používateľovi s danými právami aj upravovať výsledky nášho algoritmu a upresniť, či inak preorganizovať rozdelenie dokumentu na články, pretože vplyvom nekonzistencie tlače, a aj iných faktorov, nie je možné garantovať 100 %-nú úspešnosť rozpoznania článkov z týchto údajov.

Prínos tohto projektu

Výsledná aplikácia nášho projektu pomáha urýchliť a skvalitniť proces spracovania archivovaných periodických dokumentov na úroveň analytického rozpisu periodík a tým ich sprístupniť širokému spektru používateľov. Naše riešenie prináša metódu, ako automatizovať proces extrakcie konkrétnych článkov obsiahnutých v dokumente a zároveň pre každý z nich generovať plnohodnotný bibliografický záznam s väzbou na plný text článku a jeho vizualizovaný obraz. Vďaka výsledkom nášho projektu bude možné značne urýchliť a tiež následne skvalitniť spracovanie periodických dokumentov až na úroveň analytického roz-



Obr. 5

pisu periodík v tlačenej podobe. Toto následne vedie k novým možnostiam, ako aj ku skvalitneniu sprístupňovania týchto dokumentov a rovnako aj k novým možnostiam vyhľadávania, získavania a viacnásobného použitia týchto dokumentov. Kvalitný bibliografický popis týchto dokumentov môže viesť tiež k odhaľovaniu nových poznatkov, ktoré tieto dokumenty v sebe ukrývajú či už priamo, alebo v prepojení na iné dokumenty. Bežní používatelia tak získajú nový plnohodnotný zdroj širokého spektra informácií, v ktorých sa bude môcť štandardným spôsobom vyhľadávať a vybrané dokumenty aj priamo či sprostredkované získavať. Dôležitý je fakt, že používateľ nebude nútený čítať, či inak vizuálne si spracovávať množstvo dokumentov a dlhé texty, ktoré teraz získava pri plnotextovom vyhľadávaní.

Realizácia tohto projektu je tiež krásnou ukážkou toho, že spoluprácou školy, inštitúcií štátneho či verejného sektora, spolu so súkromným sektorom je možné dosiahnuť pridanú hodnotu a výsledky, ktoré sú prospešné širokému spektru používateľov. Je to krásny príklad toho, ako nasmerovať aktivity pri spracovaní a sprístupňovaní kultúrneho dedičstva.

Možno by bolo vhodné sa nad tým zamyslieť, ako takúto spoluprácu naštartovať a následne aj aplikovať do bežnej praxe škôl a inštitúcií štátneho a verejného sektora, pretože naša krajina má unikátne kultúrne dedičstvo, ktoré predstavuje priam nevyčísľiteľné bohatstvo. Správne formy jeho sprístupňovania, ako aj využívania môžu značne prispieť aj k rozvoju a ekonomickému rastu celej krajiny.

Bc. Martin Vaško

vaskom1@stuba.sk

Ing. Andrejčíková Nadežda, PhD.

nadezda.andrejcikova@stuba.sk

(Katedra knižničnej a informačnej vedy, Filozofická fakulta, Univerzita Komenského v Bratislave)