

JAKUB NIEDBALSKI

WPROWADZENIE DO KOMPUTEROWEJ ANALIZY DANYCH JAKOŚCIOWYCH

Przykłady bezpłatnego
oprogramowania
CAQDAS

K O M P U T E R O W A A N A L I Z A D A N Y C H



WPROWADZENIE DO KOMPUTEROWEJ ANALIZY DANYCH JAKOŚCIOWYCH

Przykłady bezpłatnego
oprogramowania
CAQDAS



WYDAWNICTWO
UNIWERSYTETU
ŁÓDZKIEGO

JAKUB NIEDBALSKI

WPROWADZENIE DO KOMPUTEROWEJ ANALIZY DANYCH JAKOŚCIOWYCH

Przykłady bezpłatnego
oprogramowania
CAQDAS

K O M P U T E R O W A A N A L I Z A D A N Y C H

Jakub Niedbalski – Uniwersytet Łódzki, Wydział Ekonomiczno-Socjologiczny
Katedra Socjologii Organizacji i Zarządzania, 90-214 Łódź, ul. Rewolucji 1905 r. nr 41/43

KOORDYNATOR SERII

Jakub Niedbalski

RECENZENT

Dariusz Kubinowski

REDAKTOR INICJUJĄCY

Iwona Gos

OPRACOWANIE REDAKCYJNE

Bogusława Kwiatkowska

SKŁAD I ŁAMANIE

Munda – Maciej Torz

PROJEKT OKŁADKI

Katarzyna Turkowska

Zdjęcie wykorzystane na okładce: © Depositphotos.com/grandfailure

© Copyright by Jakub Niedbalski, Łódź 2018

© Copyright for this edition by Uniwersytet Łódzki, Łódź 2018

Publikacja jest udostępniona na licencji Creative Commons
Uznanie autorstwa-Użycie niekomercyjne-Bez utworów zależnych 4.0 (CC BY-NC-ND)

Wydane przez Wydawnictwo Uniwersytetu Łódzkiego

Wydanie I. W.08349.17.0.M

Ark. wyd. 10,0; ark. druk. 13,5

ISBN 978-83-8142-040-2

e-ISBN 978-83-8142-041-9

<https://doi.org/10.18778/8142-040-2>

Wydawnictwo Uniwersytetu Łódzkiego

90-131 Łódź, ul. Lindleya 8

www.wydawnictwo.uni.lodz.pl

e-mail: ksiegarnia@uni.lodz.pl

tel. (42) 665 58 63

Spis treści

Wprowadzenie	9
1. OPENCODE 4 – oprogramowanie wspomagające proces przeszukiwania i kodowania danych tekstowych	15
1.1. Wygląd okna programu	15
1.2. Tworzenie projektu	17
1.3. Importowanie danych	18
1.4. Analiza zaimportowanych materiałów tekstowych	19
1.4.1. Tworzenie opisów w ramach funkcji Text 2	19
1.4.2. Kodowanie danych	23
1.4.3. Przeglądanie listy kodów	28
1.4.4. Przeszukiwanie kodów	29
1.4.5. Proces kategoryzowania – tworzenie i zarządzanie syntezami pierwszego stopnia	31
1.4.6. Przeszukiwanie syntez pod względem przypisanych im kodów	33
1.4.7. Tworzenie i zarządzanie syntezą drugiego stopnia	34
1.4.8. Tworzenie memo	38
1.5. Wyszukiwanie danych w tekście	40
1.6. Funkcje przeglądania, eksportowania i drukowania danych	40
1.7. Uwagi końcowe	45
2. CATMA – profesjonalne narzędzie do analizy tekstu	49
2.1. Rozpoczynanie pracy w środowisku programu CATMA	49
2.1.1. Dodawanie tekstów	50
2.1.2. Eksportowanie tekstów	54
2.2. Tworzenie adnotacji	54
2.2.1. Dzielenie się adnotacjami	55
2.2.2. Eksportowanie i importowanie kolekcji adnotacji	55
2.3. Tworzenie korpusu dokumentów	56
2.4. Znaczniki (tagi)	57
2.4.1. Tworzenie Biblioteki znaczników	57
2.4.2. Tworzenie znaczników	58
2.4.3. Oznaczanie tekstu tagami (tagowanie)	60
2.4.4. Edycja i usuwanie znaczników	63
2.5. Właściwości	65
2.5.1. Tworzenie właściwości	65
2.5.2. Praca z właściwościami	66

3.4. Walidacja zakodowanych danych	117
3.5. Raporty i notatki	119
3.6. Zarządzanie informacjami o koncie	120
3.7. Uwagi końcowe	120
4. RQDA – narzędzie wspomagające proces analizy danych tekstowych.	123
4.1. Rozpoczynanie pracy w programie RQDA	123
4.2. Działania wykonywane na plikach (dokumentach)	125
4.3. Tworzenie listy kodów i kodowanie danych.	129
4.4. Pisanie not i sporządzanie notatek w programie RQDA	134
4.5. Systematyzacja i zarządzanie plikami.	134
4.6. Porządkowanie i zarządzanie kodami.	137
4.7. Przypadki	139
4.8. Atrybuty	140
4.9. Wykonywanie zapytań przy użyciu SQL	143
4.10. Ustawienia i praktyczne porady	144
4.11. Uwagi końcowe	145
5. STORIES MATTER – program do porządkowania i archiwizowania materiałów audio-wizualnych	147
5.1. Wygląd okna programu.	147
5.2. Rozpoczynanie pracy w programie Stories Matter	149
5.3. Tworzenie i edytowanie informacji o rozmówcach.	153
5.3.1. Podstawowe czynności związane z dodawaniem, edytowaniem i usuwaniem rozmówców	153
5.3.2. Wprowadzanie dodatkowych informacji o rozmówcach	156
5.3.3. Funkcja powielania i przenoszenia wywiadów w bazie danych	159
5.4. Tworzenie i edycja sesji.	161
5.4.1. Podstawowe czynności związane z dodawaniem, edytowaniem i usuwaniem sesji.	161
5.4.2. Przeglądanie i wprowadzanie dodatkowych informacji o sesjach.	162
5.4.3. Funkcja powielania i przenoszenia wywiadów w bazie danych	166
5.5. Tworzenie, edycja i eksportowanie klipów	168
5.5.1. Podstawowe czynności związane z tworzeniem, edytowaniem i eksportowaniem klipów	168
5.6. Funkcje on-line dostępne w programie Stories Matter	170
5.6.1. Uzyskiwanie dostępu on-line	171
5.6.2. Narzędzie scalania bazy danych.	171
5.7. Dodatkowe funkcje programu Stories Matter	173
5.7.1. Listy odtwarzania	173
5.7.2. Chmura Tagów.	175
5.7.3. Dodatkowe narzędzia bazy danych	176
5.7.4. Opcje przeszukiwania	176
5.8. Uwagi końcowe	179

6. ELAN – oprogramowanie do analizy jakościowej materiałów audiowizualnych	181
6.1. Kluczowe elementy projektu i ich definicje	181
6.1.1. Adnotacje	181
6.1.2. Poziomy	181
6.1.3. Rodzaje (typy) poziomów	182
6.1.4. Szablony	182
6.1.5. Plik adnotacji (*.eaf)	182
6.1.6. Plik multimedialny (*.mpg, *.wav, itd.)	182
6.2. Wewnętrzna organizacja plików w programie	182
6.3. Okno programu ELAN	183
6.4. Rozpoczynanie pracy w programie ELAN	185
6.5. Podstawowe funkcje programu	186
6.6. Funkcje zaawansowane	190
6.6.1. Funkcje wyświetlania warstw (poziomów)	192
6.6.2. Tworzenie, modyfikowanie i usuwanie warstw	193
6.6.3. Tworzenie, modyfikowanie i usuwanie typów warstw	195
6.6.4. Struktura hierarchiczna warstw	197
6.6.5. Tworzenie i modyfikowanie słowników	198
6.6.6. Tworzenie i edytowanie szablonów	200
6.6.7. Praca z adnotacjami	202
6.6.7.1. Dodawanie adnotacji	202
6.6.7.2. Edycja i usuwanie adnotacji	204
6.6.7.3. Dzielenie adnotacji	206
6.7. Uwagi końcowe	206
Zakończenie	209
Bibliografia	211
Introduction to computer analysis of qualitative data. Examples of free CAQDA software.	
Summary	215

Wprowadzenie

Prowadząc zajęcia ze studentami, doktorantami czy uczestnicząc w warsztatach metodologicznych, często słyszę powtarzające się pytanie: Czy i dlaczego warto korzystać z pomocy wspomaganej komputerowo analizy danych jakościowych? Na tak postawione pytanie mogę po licznych doświadczeniach odpowiedzieć, że wprowadzie CAQDA¹ jest czasochłonna, bowiem posługiwanie się programami wymaga pewnego wysiłku i nakładów pracy związanych z poznaniem środowiska danego oprogramowania, ale także (a może przede wszystkim) zmiany optyki, a nierzadko utartych przyzwyczajęń dotyczących organizacji warsztatu badacza (Lofland, Snow, Anderson, Lofland 2009: 144–145; Kubinowski 2010: 73–77), jednak korzystanie z CAQDA sprawia, że proces analizy danych jest bardziej systematyczny i przejrzysty. To niejako „zmusza” do poważnego myślenia o danych oraz sprzyja budowaniu teorii. Co więcej, w zależności od paradygmatu metodologicznego można dzięki stosowaniu CAQDA przeprowadzić jakościową analizę porównawczą bądź zbadać kilka hipotez naraz (zob. Kelle 1995).

Dlatego zasadniczym celem, który postawiłem sobie jako autor niniejszej książki, jest próba zaprezentowania z perspektywy badacza jakościowego, a zarazem użytkownika oprogramowania CAQDAS², możliwości oraz sposobów korzystania z wybranych programów wspomagających analizę danych jakościowych. Niniejsza monografia, podobnie jak jej poprzedniczka (Niedbalski 2013a), skierowana jest przede wszystkim do badaczy, którzy dopiero chcą spróbować swoich sił w realizacji własnych projektów badawczych z wykorzystaniem programów z rodziny CAQDAS. W ich przypadku dobrym pomysłem może okazać się bezpłatne oprogramowanie dystrybuowane na zasadach wolnej licencji. Choć programy te są zazwyczaj uboższe w porównaniu do wersji płatnych, nie posiadają bowiem tak rozbudowanych funkcji lub oferują mniej przyjazne środowisko pracy dla użytkownika, to w większości stanowią dobrą alternatywę dla nadal dość drogich i z tego względu nie zawsze powszechnie dostępnych programów odpłatnych. Także i tym razem, jak miało to miejsce podczas wydawania pierwszej mojej książki poświęconej CAQDAS (Niedbalski 2013a), zdecydowałem się

¹ Posługując się w książce skrótem CAQDA, a więc Computer Assisted Qualitative Data Analysis, mam na myśli wspomaganą komputerowo analizę danych jakościowych.

² Posługując się w książce skrótem CAQDAS, a więc Computer Assisted Qualitative Data Analysis Software, mam na myśli oprogramowanie służące do komputerowego wspomaganie analizy danych jakościowych.

zaprezentować narzędzia bezpłatne, powszechnie dostępne za pośrednictwem stron internetowych, a zarazem w pełni funkcjonalne, tj. nieposiadające żadnych ograniczeń ze strony ich wydawców czy autorów. Uznałem bowiem, że w ten sposób użytkownik otrzymuje program, który bez żadnych dodatkowych ograniczeń może wykorzystywać w swoich badaniach.

Ponadto opisane programy stanowią swego rodzaju przekrój możliwości oprogramowania CAQDAS. Moją intencją było bowiem przedstawić kilka narzędzi, które nie będą reprezentowały identycznego zestawu funkcji. Przy czym w niniejszym opracowaniu duży nacisk położyłem nie tylko na programy pozwalające na pracę z tekstem, ale też te, które umożliwiają analizę materiałów audiowizualnych. Mam nadzieję, że dzięki temu badacze jakościowi reprezentujący różne szkoły i posługujący się odmiennymi metodologiami znajdą narzędzia odpowiednie dla siebie.

Przy wyborze konkretnych programów kierowałem się także walorami użytkowymi, związanymi ze specyfiką środowiska pracy oraz ich funkcjonalnością. Z tego względu wszystkie opisane w książce narzędzia cechuje jasna budowa, klarowna architektura i intuicyjne rozmieszczenie poszczególnych opcji. Dzięki temu praktycznie każda osoba, która posiada nawet podstawowe kompetencje w zakresie obsługi popularnych programów, służących m.in. do edycji tekstu, nie powinna mieć w ich wypadku większych trudności.

Jednocześnie, oprócz stosunkowo łatwej obsługi opisanego oprogramowania istotny dla jego wyboru był sam proces instalacji. Nie widzę bowiem powodu, aby wymagać od przeciętnego użytkownika posiadania wiedzy z zakresu informatyki. Z tego względu ów proces w przypadku opisanych programów ogranicza się właściwie do ich ściągnięcia z odpowiedniej strony internetowej, a następnie uruchomienia i postępowania zgodnie z wyświetlanymi instrukcjami, zazwyczaj ograniczającymi się do wyboru bądź akceptacji kolejnych kroków. Co więcej, wychodząc naprzeciw przyszłym użytkownikom, podjąłem decyzję o wyborze narzędzi dostępnych on-line, a więc takich, które nie wymagają procesu instalacji.

Wreszcie dokonany wybór jest sumą moich osobistych doświadczeń, o które na przestrzeni ostatnich kilku lat jestem też bogatszy. Przy czym zdaję sobie sprawę, że są to kwestie wysoce indywidualne i z tego względu rozumiem, że zdania mogą być podzielone, a oczekiwania poszczególnych użytkowników – zróżnicowane. Dlatego w dalszym ciągu podtrzymuję moje wcześniejsze słowa, że nikogo nie przekonuję na siłę do korzystania właśnie z tych programów (czy w ogóle z oprogramowania CAQDAS). Pragnę jedynie wskazać na ich możliwości i pokazać „techniczne” ułatwienia w zakresie pracy analitycznej, wiążące się z ich wykorzystaniem.

Książka została poświęcona przedstawieniu pięciu programów CAQDAS, które stanowią swego rodzaju przekrój dostępnych obecnie narzędzi wspomagających analizę danych jakościowych. Skupiłem się głównie na tych programach,

które oferują dość bogate możliwości, są bezpłatne, powszechnie dostępne oraz nie wymagają specjalnych zasobów sprzętowych ani wiedzy z zakresu obsługi standardowych programów komputerowych (takich jak na przykład MS Office). W związku z tym w niniejszym opracowaniu nie zabrakło przykładów oprogramowania, które pozwalają zarówno na pracę z tekstem, jak i analizę danych audiowizualnych, a także dostępnych on-line, bez konieczności ich bezpośredniej instalacji na komputerze użytkownika.

Pierwszym z programów, które zamieściłem w książce, jest znany już z mojego wcześniejszego opracowania (Niedbalski 2013a) **OPENCODE**. Przy czym wracam do niego celowo, ponieważ w ciągu kilku ostatnich lat przeszedł on dość gruntowne modyfikacje, przez co zmieniła się jego specyfika oraz zakres możliwego zastosowania. W dalszym ciągu jest to przede wszystkim narzędzie do kodowania danych jakościowych występujących w formie tekstowej, takich jak wywiady czy obserwacje. Nadal też są w nim wyraźnie widoczne nawiązania do procedur i zasad znanych z metodologii teorii ugruntowanej (Glaser, Strauss 1967 [2009]; Glaser 1978; Strauss, Corbin 1990; Konecki 2000; Gorzko 2008; Charmaz 2006 [2009]). Niemniej jednak po ostatnich ulepszeniach zyskał on dodatkową funkcję, umożliwiającą analizę treści, a dzięki zastosowaniu rozbudowanej formuły syntetyzowania danych i tworzenia drzewa kategorii znacznie wzrosły jego zalety jako narzędzia do klasyfikacji oraz sortowania wszelkiego rodzaju informacji tekstowych, których analiza jest prowadzona w duchu metod jakościowych (zob. Knoblauch, Flick, Maeder 2005; Flick 2010; Rapley 2010).

Kolejny program, **CATMA**, służy do analizy danych tekstowych. Pozwala na zestawienia dokumentów w korpusie, tworząc ich kolekcję. Co istotne, informacje zawarte w programie mogą być współdzielone z innymi osobami, co możliwe jest dzięki temu, że jest on implementowany jako aplikacja internetowa. Wśród innych ważnych cech oprogramowania wymienić można: obsługę cyfrowego tekstu w niemalże każdym języku, dowolnie definiowalne lub predefiniowane znaczniki, interaktywne wyszukiwanie w języku naturalnym poprzez tekst, korpus i adnotacje, a także automatyczne, statystyczne i niestatystyczne funkcje analityczne czy wbudowaną wizualizację wyników wyszukiwania i przeprowadzanych analiz.

Również **CAT**, a więc Coding Analysis Toolkit opisany jako trzeci z kolei program w niniejszej książce, jest internetowym zestawem narzędzi CAQDAS. CAT umożliwia wydajną, przejrzystą, niezawodną, prawidłową i skalowalną, opartą na sieci współpracę dotyczącą kodowania i analizy tekstów. Program umożliwia zarządzanie danymi, kodami, a także analizę materiałów tekstowych. Posiada również rozbudowany system raportowania, dzięki czemu możliwe jest generowanie informacji będących rezultatem prowadzonej pracy analitycznej. Do niewątpliwych atutów programu należą łatwość obsługi, intuicyjność oraz przejrzysty interfejs, a jako system oparty na sieci internetowej stwarza okazję do szybkiego

i wygodnego replikowania oraz łączenia danych, jak również sprawdzania wykonanych analiz przez zespół badaczy.

Kolejnym narzędziem, któremu poświęciłem następny rozdział w książce, jest służący do analizy tekstu program **RQDA**. W istocie RQDA jest służącym do analizy danych jakościowych pakietem dobrze znanego w środowisku badaczy ilościowych oprogramowania R. Oba programy tworzą zresztą zintegrowaną platformę, która doskonale wpisuje się w klimat badań mieszanych: jakościowo-ilościowych lub ilościowo-jakościowych. Wśród wielu opcji, w jakie wyposażony został program RQDA, które wspomagają proces analizy danych, można wyróżnić: kodowanie, segregowanie i porządkownie wszystkich zaimportowanych oraz utworzonych danych, przeszukiwanie zakodowanych informacji, w tym z zastosowaniem operatorów logicznych, wyszukiwanie w plikach słów kluczowych, tworzenie przypadków oraz atrybutów, a także wszechstronną i elastyczną możliwość edytowania i organizowania poszczególnych plików.

STORIES MATTER to kolejny program, którego opis zamieściłem w niniejszym opracowaniu. Jest to narzędzie stworzone z myślą o archiwizacji cyfrowych materiałów wideo i audio. Zgodnie z zamysłem jego twórców, ma on stanowić alternatywę dla transkrypcji, a więc umożliwiać gromadzenie, ale także sortowanie i porządkowanie materiałów. Program pozwala na edytowanie nagrań według własnych kryteriów użytkownika, a także na tworzenie zestawień zawierających pliki audio i wideo uporządkowane według określonych tematów czy problemów badawczych. Natomiast dzięki systemowi notatek oraz opisów **STORIES MATTER** pozwala na wprowadzanie dodatkowych informacji o gromadzonych plikach i ich zawartości, a system przeszukiwania oraz tworzenia chmury tagów zdecydowanie usprawnia proces rozpoznawania kluczowych tematów i kwestii poruszanych przez rozmówców. Użytkownicy mogą również eksportować wyniki swoich prac w kilku różnych formatach, co ułatwia ich wykorzystanie w prezentacjach czy przy projektowaniu witryn internetowych.

Ostatni z prezentowanych programów **ELAN** to program stworzony z myślą o osobach, które chcą wykonać analizę danych wizualnych z zastosowaniem narzędzi komputerowych. **ELAN** pozwala tworzyć, edytować, wizualizować i wyszukiwać adnotacje w przypadku danych wideo i audio. Jest to narzędzie przeznaczone specjalnie do analizy języka, w tym także do języka migowego i gestów, ale ma charakter uniwersalny i może być z powodzeniem używane przez wszystkich, którzy pracują z danymi wideo i/lub audio w celach analizy oraz dokumentacji danych. Niewątpliwym atutem programu jest czytelny interfejs, a także przyjazny dla użytkownika layout. Z tego względu **ELAN** to dobre narzędzie dla tych wszystkich, którzy poszukują alternatywy narzędzi **CAQDAS** dla drogich programów licencjonowanych, a chcą wykonać w sposób profesjonalny analizę jakościową opartą na danych audiowizualnych (zob. Konecki 2012, 2009).

Podsumowując, chciałbym jeszcze raz podkreślić, iż moim zamiarem było przedstawić programy, których zadaniem jest wspomaganie pracy badacza jakościowego, reprezentującego różne szkoły i wykorzystującego rozmaite metody analityczne. Trzeba jednak pamiętać, że każdy program jest swoistym „środowiskiem”, w którym badacz pracuje i wykonuje określone czynności zgodnie z tzw. „architekturą oprogramowania”, a więc technicznymi rozwiązaniami użytymi przez jego konstruktorów (Saillard 2011: 2). Przy czym twórcy tego rodzaju oprogramowania podejmują starania, aby nie nakładały one ograniczeń natury metodologicznej. To zaś, który z omówionych tutaj programów ostatecznie znajdzie uznanie u danego użytkownika, jest w istocie sprawą indywidualną. Wiele zależy bowiem od tego, jakie są potrzeby konkretnego badacza, a to z kolei uwarunkowane jest zarówno stosowanymi przez niego metodami, jak i podejmowaną problematyką oraz osobistymi preferencjami (Lonkila 1995; Saillard 2011: 3).

Przede wszystkim jednak – na co zawsze staram się zwracać uwagę zarówno czytelników, jak i osób uczestniczących w prowadzonych przeze mnie warsztatach – kluczowe jest, aby w pełni zdawać sobie sprawę ze **wspomagającego** charakteru używanego oprogramowania. Musimy bowiem pamiętać, że żadne, nawet najbardziej wyrafinowane i zaawansowane technologicznie programy nie wyręczą badacza (Dohan, Sanchez-Jankowski 1998: 482; Bringer, Johnston, Brackenridge 2004: 249). To od refleksyjności, posiadanej wiedzy i doświadczenia badacza będą zależały wyniki podejmowanych przez niego działań (Lonkila 1995; Silverman 2008: 101). Innymi słowy, jedynym odpowiedzialnym za poziom analizy i jakość wykonanej pracy nieodłącznie pozostaje sam badacz (Coffey, Atkinson 1996: 187; Bringer, Johnston, Brackenridge 2006: 247). W związku z tym nie należy tego typu oprogramowania traktować jako swoistego remedium na problemy konceptualne czy trudności związane z interpretacją danych, a oczekiwane rezultaty można jedynie osiągnąć, łącząc dwie podstawowe role: świadomego badacza-analityka, a zarazem wprawnego użytkownika danego programu (por. Miles, Huberman 2000).

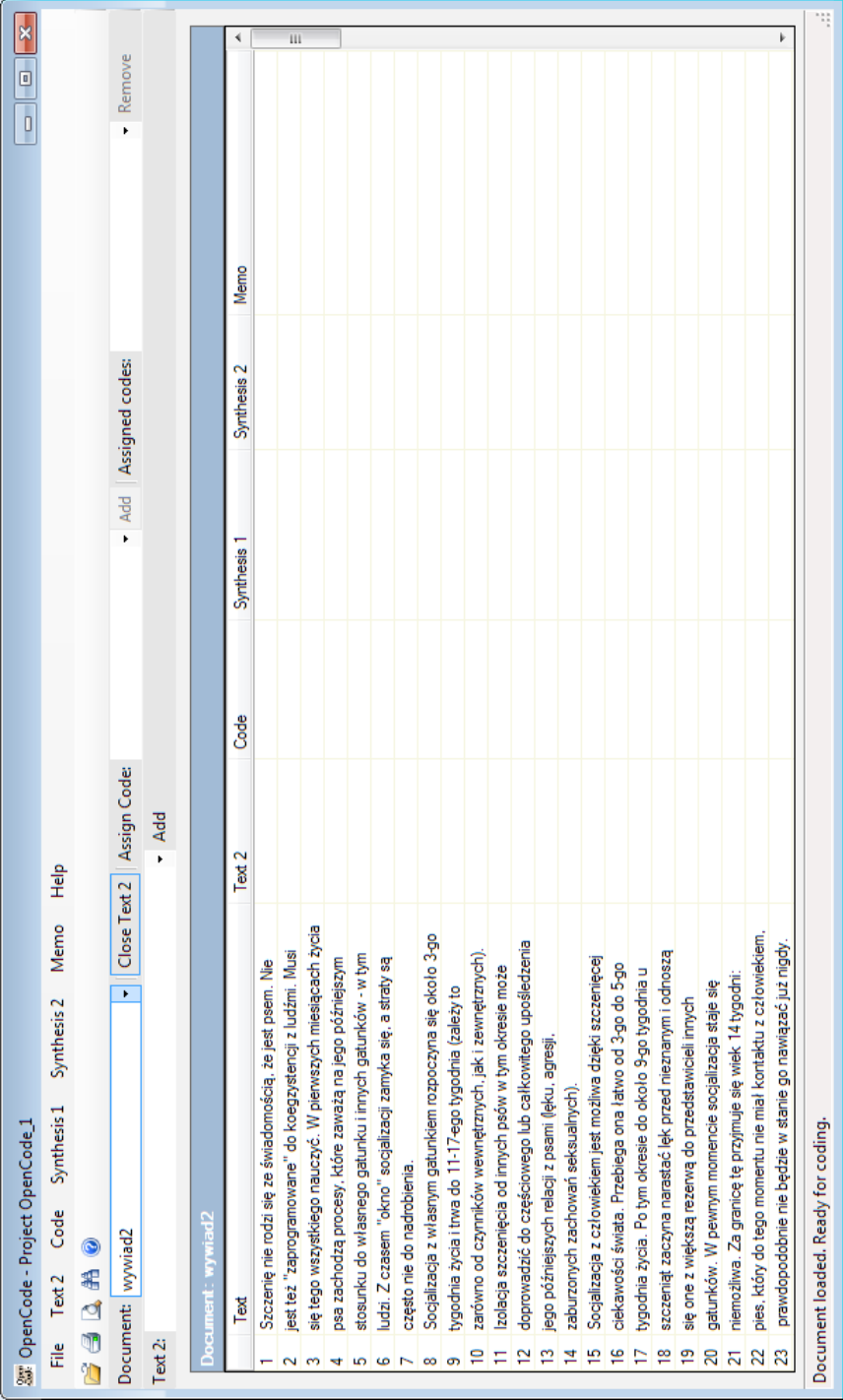
1. OPENCODE – oprogramowanie wspomagające proces przeszukiwania i kodowania danych tekstowych¹

OpenCode jest oferowany bezpłatnie i dostępny na stronie Departamentu Zdrowia Publicznego i Medycyny Klinicznej Szwedzkiego Uniwersytetu w Umeå, a także dystrybuowany razem z książką Larsa Dahlgrena, Marii Emmelin i Anny Winkvist (2007) pod tytułem „Qualitative methodology for international public health”. Jak podają autorzy programu, ich celem było stworzenie narzędzia wspomagającego jakościową analizę materiałów tekstowych, które jednocześnie jest łatwe do opanowania przez przeciętnego użytkownika i proste w obsłudze (Niedbalski 2012: 221). Program OpenCode to narzędzie do kodowania danych jakościowych występujących w formie tekstowej (przy czym obsługiwany format to .txt), takich jak wywiady czy obserwacje. Został on opracowany na potrzeby analizy prowadzonej zgodnie z zasadami metodologii teorii ugruntowanej (Glaser, Strauss 1967; Glaser 1978; Konecki 2000; Charmaz 2009). Niemniej jednak może być z powodzeniem używany jako narzędzie do klasyfikacji i sortowania wszelkiego rodzaju informacji tekstowych, których analiza prowadzona jest w duchu metod jakościowych. OpenCode jako narzędzie wspomagające proces analizy danych jakościowych posiada dość rozbudowane możliwości. Spośród dostępnych funkcji należy przede wszystkim wskazać na możliwość: tworzenia bazy danych materiałów tekstowych, przeszukiwania tekstów pod kątem określonych słów, przypisywania kodów do określonych segmentów tekstu, tworzenia i zarządzania kategoriami służącymi do grupowania wygenerowanych kodów, przeglądania i przeszukiwania utworzonych kodów oraz kategorii czy tworzenia notatek w formie memo do zapisywania krótkich informacji bądź myśli analitycznych badacza.

1.1. Wygląd okna programu

Program posiada bardzo prosty i intuicyjny wygląd. Autorzy programu zadbali, aby wszystkie istotne funkcje były dostępne „na wyciągnięcie ręki”. W głównym oknie programu u samej góry znajduje się menu główne, gdzie można uzyskać

¹ Rozdział jest zmienioną i unowocześnioną wersją tekstu zamieszczonego w książce mojego autorstwa, która ukazała się w 2013 r. pt. *Odkrywanie CAQDAS. Wybrane bezpłatne programy komputerowe wspomagające analizę danych jakościowych*, Wydawnictwo UŁ.



Ilustracja 1.1. Wygląd okna programu OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

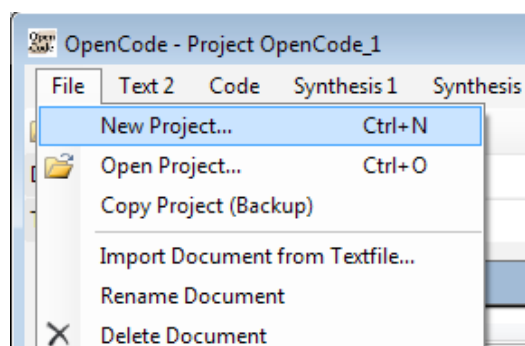
dostęp do większości funkcji, w tym: kodowania, kategoryzowania oraz pisania memo. Nieco niżej mamy prosty pasek narzędzi w formie ikon, które umożliwiają szybki dostęp do wybranych funkcji, takich jak otworenie projektu, opcje drukowania, podgląd wydruku, przeszukiwanie oraz pomoc. Jeszcze niżej znajduje się pasek narzędzi, gdzie można wybrać dokument (np. wywiad bądź notatkę z obserwacji), który badacz zamierza aktualnie opracowywać. Obok, w tym samym miejscu znajduje się także podstawowa z punktu widzenia analizy metodologii teorii ugruntowanej opcja przypisywania i usuwania kodów dla wybranych linii (wersów) dokumentu.

W głównej części okna wyświetlany jest dokument w formie tabeli. Pierwsze dwie kolumny zawierają numer linii oraz tekst dokumentu. Następne kolumny prezentują kolejno: tak zwany text 2, przypisane segmentom tekstu kody oraz nazwy istniejących syntez pierwszego i drugiego stopnia (dawniej kategorii), a także nazwy memo.

Tak jak intuicyjny jest układ menu i wygląd całego programu, tak samo prosty jest sposób wykonywania poszczególnych czynności analitycznych za pomocą funkcji dostępnych w programie.

1.2. Tworzenie projektu

Pierwszą czynnością, jaką należy wykonać w programie, jest utworzenie projektu, czyli bazy danych, do której będą następnie importowane i przechowywane kolejne dokumenty tekstowe, np. transkrypcje wywiadów. W tym celu należy z menu *File* (Plik) wybrać opcję *New Project...* (Nowy Projekt), a następnie podać nazwę projektu. Domyślnie każdy nowy projekt zapisywany jest w katalogu *My OpenCode Project* (Moje Projekty) znajdującym się w katalogu *Moje dokumenty*.



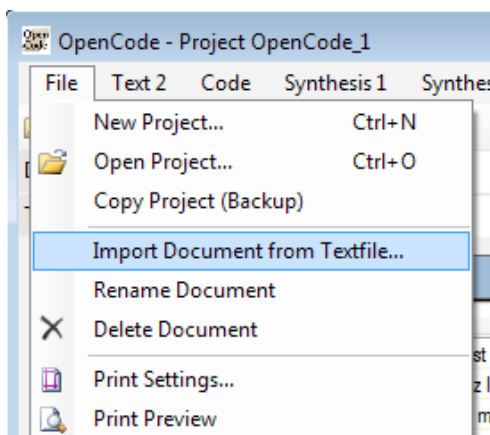
Ilustracja 1.2. Tworzenie projektu w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

Po czynności utworzenia projektu możemy już przystąpić do gromadzenia dokumentów. Należy przy tym pamiętać, że program OpenCode nie obsługuje innych formatów niż .txt. Z tego względu przed zaimportowaniem dokumentu należy go zapisać jako „zwykły tekst” w jednym z dostępnych programów do edycji tekstów.

1.3. Importowanie danych

Zaimportowany dokument zostaje automatycznie podzielony na ponumerowane wersy, każdy o długości sześćdziesięciu znaków. Pewną niedogodność może sprawiać fakt, że po zaimportowaniu dokumentu nie ma możliwości jego formatowania i edytowania, dlatego wszelkich korekt należy dokonywać przed umieszczeniem tekstu w bazie danych. Z tego względu, jeżeli zaimportowany plik z jakiś przyczyn okaże się nieprawidłowy, na przykład będą za duże odstęp między wersami, zamiast tekstu pojawią się „szlaczki” lub po prostu po wnikliwym przejrzaniu tekstu zostaną wykryte innego rodzaju błędy m.in. „literówki”, wówczas najlepszym rozwiązaniem jest usunąć plik z bazy (*File/Delete Document*) i zaimportować go ponownie. Dlatego zaleca się przed rozpoczęciem pracy nad dokumentem dokładne jego sprawdzenie pod kątem wspomnianych usterek.



Ilustracja 1.3. Importowanie dokumentu tekstowego w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

Po zaimportowaniu materiału (*File/Import Document from Textfile...*) badacz może przystąpić do jego opracowywania. Pierwszą czynnością, o ile nie uczyniono tego w momencie importowania danych, powinno być ich opisanie. W tym

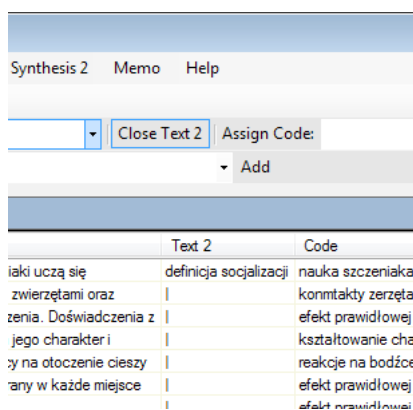
celu należy skorzystać z opcji *Document Information* (Informacje o dokumencie) znajdującej się w menu *File* (Plik). W otwartym oknie badacz wpisuje informacje odnoszące się do określonego dokumentu, na przykład warunki i okoliczności przeprowadzania wywiadu czy dane metryczkowe rozmówcy.

Warto przy tym dodać, że użytkownik w każdej chwili ma możliwość edycji danych zawartych w opisie dokumentu, a także zmiany jego nazwy. Jest to szczególnie ważne, jeśli weźmiemy od uwagę fakt, że w toku analizy danych badacz może chcieć uzupełniać informacje dotyczące analizowanych wywiadów czy notatek terenowych, ale też zmieniać nazwy plików, jeśli te mają za zadanie porządkować materiały w miarę postępowania procesu analitycznego.

1.4. Analiza zaimportowanych materiałów tekstowych

1.4.1. Tworzenie opisów w ramach funkcji Text 2

Jak sugerują sami twórcy programu, we wczesnych etapach analizy danych użyteczna może się okazać funkcja Text 2, która pozwala na tworzenie swego rodzaju podsumowań czy też opisów dotyczących poszczególnych segmentów tekstu pierwotnego. To także opcja przydatna, gdy pracujemy nad jakościową analizą treści.



The screenshot shows the 'Synthesis 2' application window with a menu bar (Synthesis 2, Memo, Help) and a toolbar with buttons like 'Close Text 2' and 'Assign Code'. Below the toolbar is a table with two columns: 'Text 2' and 'Code'. The table contains several rows of text segments and their corresponding codes.

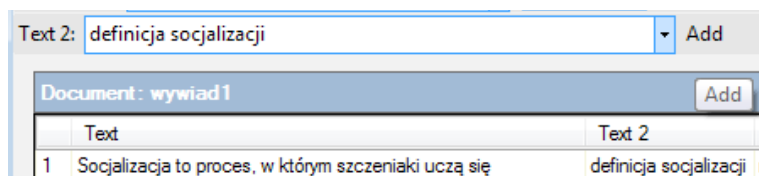
	Text 2	Code
	iaiki uczą się	definicja socjalizacji
	zwierzętami oraz	nauka szczeniaka,
	zenia. Doświadczenia z	kontakty zerkętar
	jego charakter i	efekt prawidłowej s
	zy na otoczenie cieszy	kształtowanie char
	any w każde miejsce	reakcje na bodźce
		efekt prawidłowej s
		efekt prawidłowej s

Ilustracja 1.4. Kolumna Text 2 w obszarze roboczym okna programu OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

Ponieważ użycie tej funkcji nie jest obligatoryjne, a w dużej mierze zależne od potrzeb i preferencji badacza, stąd konstruktorzy programu wyszli z założenia, że będzie ona dostępna dopiero po kliknięciu na przycisk *Open Text 2* w obszarze paska narzędzi. Wykonanie tej czynności spowoduje wyświetlenie się obszaru, gdzie można dokonywać wpisów opisujących tekst pierwotny.

Aby użyć tej opcji, należy wybrać żądane wiersze, a następnie wpisać tekst w odpowiednim polu i kliknąć na przycisk *Add* (Dodaj). W ten sposób taki tekst pojawi się w kolumnie Text 2 przy wybranym wersie (lub wersach). Przy czym, jeśli co najmniej jedna z zaznaczonych linii ma już taki tekst, wówczas zostanie on zastąpiony nowym wpisem.



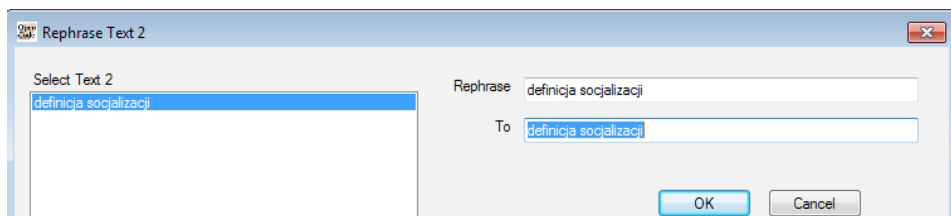
Ilustracja 1.5. Dodawanie opisu do kolumny Text 2

Źródło: opracowanie własne na podstawie programu OpenCode 4

Warto również pamiętać, że maksymalna długość tekstu nie może przekraczać 200 znaków. Ponadto, jeśli wybrano kilka linii, to wpisany tekst pojawi się tylko w pierwszym wierszu, a w kolejnych będą widoczne pionowe linie.

Jeśli z jakiś przyczyn okaże się, że będziemy chcieli dokonać zmian w istniejącym już wpisie należącym do Text 2, wystarczy, że wybierzemy z menu funkcję *Rephrase Text 2*, a następnie w wyświetlonym oknie w polu edycji wpiszę nową (lub poprawioną) wersję. Natomiast całkowite usunięcie wpisu z danej linii nastąpi wówczas, gdy zamiast tekstu pozostawimy puste pole i klikniemy na przycisk *Add* (Dodaj).

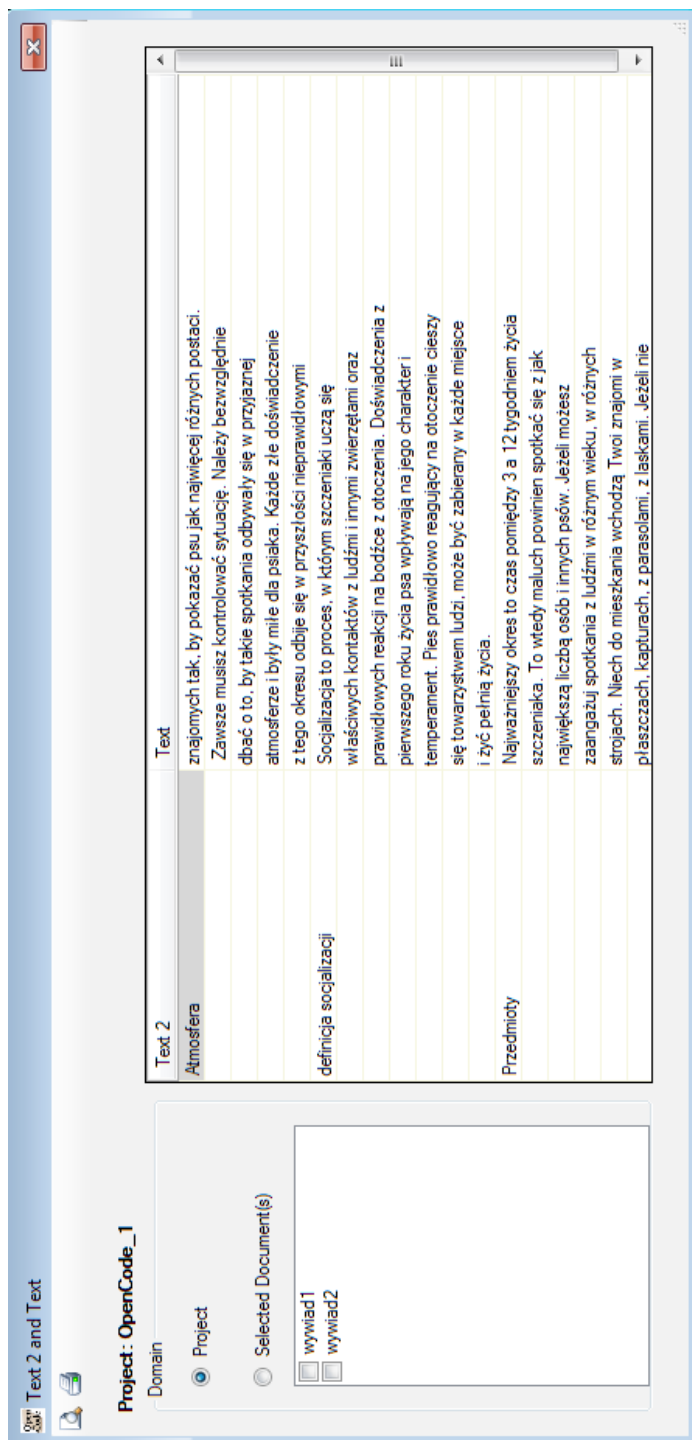
Uwaga: Wybranie opcji *Rephrase Text 2* spowoduje zmianę tekstu we wszystkich miejscach, w których został on użyty.



Ilustracja 1.6. Użycie funkcji *Rephrase Text 2*

Źródło: opracowanie własne na podstawie programu OpenCode 4

Program umożliwia także wgląd w dotychczas utworzone wpisy, dzięki systemowi zbiorczych list. Możemy zatem uzyskać pełną listę wpisów zarówno odnośnie do całego projektu, jak i wybranych dokumentów. W tym celu należy z menu *Text 2* wybrać opcję *Text 2 and Text List*.



Ilustracja 1.7. Okno z wyświetloną listą wpisów po wyborze opcji *Text 2 and Text List*

Źródło: opracowanie własne na podstawie programu OpenCode 4



Ilustracja 1.8. Okno z wyświetloną listą wpisów po wyborze opcji *Text 2 and Code List*

Źródło: opracowanie własne na podstawie programu OpenCode 4

Natomiast wybór opcji *Text 2 and Code List* spowoduje, że uzyskamy listę wszystkich wpisów w zestawieniu z kodami w ramach poszczególnych wersji, co także możemy uczynić odnośnie do całego projektu lub wybranych dokumentów.

1.4.2. Kodowanie danych

Analizę materiału znajdującego się w bazie danych rozpoczynamy od kodowania, przypisując poszczególnym segmentom tekstu określone kody. Aby wykonać tę czynność, badacz powinien zaznaczyć wers lub wersy, w których znajdują się interesujące go fragmenty tekstu, a następnie wpisać nowy kod w polu znajdującym się na pasku narzędzi (*Assign Code*) bądź wybrać jeden z istniejących już kodów widocznych w formie rozwijanej listy, a następnie w celu potwierdzenia i zapisania kodu nacisnąć *Add* (Dodaj) lub użyć skrótu klawiaturowego Ctrl + Q.

Wszystkie kody zarówno nowo utworzone, jak i przypisane spośród już istniejących są na bieżąco zapisywane do bazy danych/projektu, aby zapobiec utracie informacji podczas wychodzenia z programu.

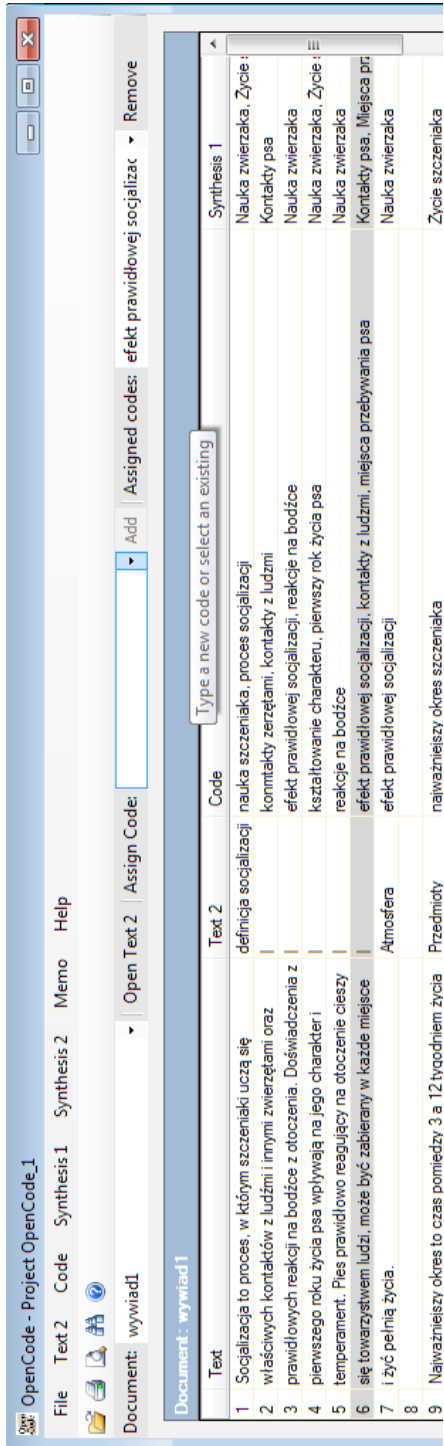
Trzeba także pamiętać, że maksymalna długość kodu wynosić może nie więcej niż 50 znaków. Nie ma natomiast ograniczeń, co do używania polskich znaków oraz dużych i małych liter. Można także używać dowolnych znaków do utworzenia nazwy kodu za wyjątkiem przecinków.

Nie ma również ograniczeń, co do liczby kodów przypisanych danemu segmentowi tekstu oraz pojedynczemu wierszowi. Nie można jednak przypisywać kodów do pustych linii, w których nie ma tekstu.

Usuwanie kodów odbywa się poprzez zaznaczenie zakodowanego segmentu tekstu, a następnie wybór konkretnego kodu w polu *Assigned Codes* (Przypisane Kody) na pasku menu i naciśnięciu opcji *Remove* (Cofnij).

Opisane funkcje dostępne są także po naciśnięciu prawego przycisku myszy i otwarciu menu kontekstowego, a następnie wybraniu odpowiedniej opcji dodania lub usunięcia kodu spośród kodów już istniejących. Przy czym należy pamiętać, że kursor powinien znajdować się w kolumnie kodów na zaznaczonym wierszu (wierszach).

Jeśli zamierzamy zmienić kod, można to uczynić na dwa sposoby, w zależności od tego, czy zmiana ma dotyczyć konkretnego segmentu tekstu, czy całego projektu. W pierwszym wypadku wybieramy linie, w których pojawia się kod, po czym usuwamy go, a w jego miejsce dodajemy nowy. W drugim wypadku, kiedy decydujemy się na zmianę kodu w całym projekcie, powinniśmy skorzystać z opcji *Rename Code* dostępnej w menu *Code*, a następnie wybrać kod, którego nazwę chcemy zmienić i w odpowiednim polu wpisać jego nową nazwę.



Ilustracja 1.9. Przypisywanie nowego kodu w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

OpenCode - Project OpenCode_1

FileText 2CodeSynthesis 1MemoHelp

Document: wywiad1

Open Text 2Assign Code: AddAssigned c

Document: wywiad 1

Text	Text 2	Code
1	Socjalizacja to proces, w którym szczeniaki uczą się	nauka szczeniaka, proces socjalizacji
2	właściwych kontaktów z ludźmi i innymi zwierzętami oraz	kontakty z zwierzętami, kontakty z ludźmi
3	prawidłowych reakcji na bodźce z otoczenia. Doświadczenia z	efekt prawidłowej socjalizacji, reakcje na bodźce
4	pierwszego roku życia psa wpływają na jego charakter i	kształtowanie charakteru, pierwszy rok życia psa
5	temperament. Pies prawidłowo reagujący na otoczenie cieszy	reakcje na bodźce
6	się towarzystwem ludzi, może być zabierany w każde miejsce	efekt prawidłowej socjalizacji, kontakty z ludźmi, miejsca pr
7	i żyć pełnią życia.	efekt p
8		
9	Najważniejszy okres to czas pomiędzy 3 a 12 tygodniem życia	najważ
10	szczeniaka. To wtedy maluch powinien spotkać się z jak	dostarc
11	największą liczbą osób i innych psów. Jeżeli możesz	kontakt
12	zaangażuj spotkania z ludźmi w różnym wieku. Jeżeli możesz	dostarc
13	strojach. Niech do mieszkania wchodzi Twój znajomy w	dostarc
14	placzkach, kapturach, z parasolami, z laskami. Jeżeli nie	dostarc
15	masz dzieci poproś o spotkanie z dziećmi sąsiadów,	dostarc
16	znajomych. Psy często źle reagują na postawnych mężczyzn,	dostarczanie bodźców, kontakty z mężczyznami, nazwa ko
17	szczególnie jeżeli domownicy to kobiety i dzieci. Dobieraj	kontakty z dziećmi, kontakty z kobietami

dobrze zachowanie w różnych sytuacjach

dostarczanie bodźców

efekt prawidłowej socjalizacji

gabinet weterynary

kontakty z zwierzętami

kontakty z dziećmi

kontakty z kobietami

kontakty z ludźmi

kontakty z mężczyznami

kontrola sytuacji

kształtowanie charakteru

lekarz weterynary

miejsca przebywania psa

najważniejszy okres szczeniaka

nauka psa

nauka szczeniaka

nazwa kodu

negatywne reakcje obcego psa

negatywne reakcje psa

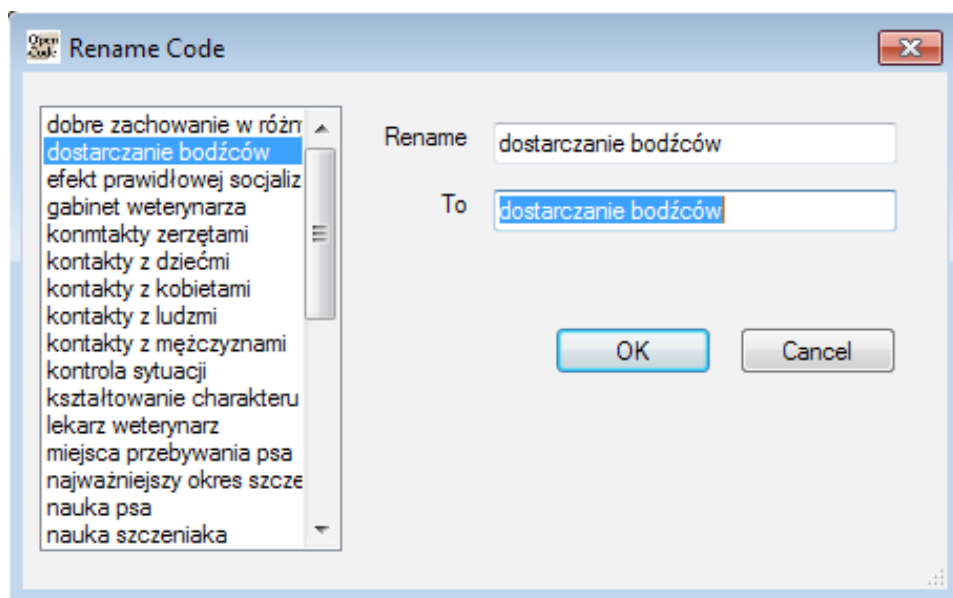
nieprawidłowa reakcja na bodźce

nieśmiałość szczeniąt

nowe sprzęty

Ilustracja 1.10. Menu kontekstowe – dodawanie/usuwanie kodów w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4



Ilustracja 1.11. Zmiana nazwy kodu w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

Uwaga: Usuwanie kodu dotyczy danego wiersza/wierszy zaznaczonych, nie zaś wszystkich miejsc w dokumencie, w których użyliśmy owego kodu. Zmiana nazwy ma natomiast wymiar generalny i odnosi się natomiast do wszystkich wierszy, do których użyliśmy określonego kodu.

Jeśli badacz chce przypisać dany kod do segmentu tekstu zajmującego więcej niż jeden wers, wówczas powinien wykonać jedną z podanych niżej czynności:

- zaznaczyć kilka kolejnych wersów – w tym celu należy kliknąć na pierwszej linii i przytrzymać przycisk myszy, a następnie przeciągnąć kursor do ostatniej wybranej linii. Ten sam efekt uzyskamy, przytrzymując klawisz Shift;
- zaznaczyć kilka wersów niewystępujących kolejno po sobie – aby to zrobić, trzeba, trzymając wciśnięty klawisz Ctrl, kliknąć kolejno na żądanych wersach;
- zaznaczyć cały dokument – w tym celu należy skorzystać z kombinacji klawiszy Ctrl + A na klawiaturze;
- odznaczyć któryś z uprzednio wybranych wersów – czynność tę można wykonać, trzymając wciśnięty klawisz Ctrl, klikając na linię, którą chcemy odznaczyć.

OpenCode - Project OpenCode_1

File

Text 2

Code

Synthesis 1

Synthesis 2

Memo

Help

Document: wywiad1

Open Text 2

Assign Code:

Add

Assigned codes: dostarczanie bodźców

Remove

Document : wywiad1

	Text	Text 2	Code	Synthesis 1
1	Socjalizacja to proces, w którym szczeniaki uczą się	definicja socjalizacji	nauka szczeniaka, proces socjalizacji	Nauka zwierzaka, Życie :
2	właściwych kontaktów z ludźmi i innymi zwierzętami oraz		konnтакты zerezłami, konтакты z ludźmi	Kontakty psa
3	prawidłowych reakcji na bodźce z otoczenia. Doświadczenia z		efekt prawidłowej socjalizacji, reakcje na bodźce	Nauka zwierzaka
4	pierwszego roku życia psa wpływają na jego charakter i		kształtowanie charakteru, pierwszy rok życia psa	Nauka zwierzaka, Życie :
5	temperament. Pies prawidłowo reagujący na otoczenie cieszy		reakcje na bodźce	Nauka zwierzaka
6	się towarzystwem ludzi, może być zabierany w każde miejsce		efekt prawidłowej socjalizacji, konтакты z ludźmi, miejsca przebywania psa	Kontakty psa, Miejsca pr:
7	i żyć pełnią życia.	Atmosfera	efekt prawidłowej socjalizacji	Nauka zwierzaka
8				
9	Najważniejszy okres to czas pomiędzy 3 a 12 tygodniem życia	Przedmioty	najważniejszy okres szczeniaka	Życie szczeniaka
10	szczeniaka. To wtedy maluch powinien spotkać się z jak		dostarczanie bodźców, nauka szczeniaka	Nauka zwierzaka, nowa :
11	większą liczbą osób i innych psów. Jeżeli możesz		konnтакты zerezłami, konтакты z ludźmi, nauka szczeniaka	Kontakty psa, Nauka zwi
12	zaangażuj spotkania z ludźmi w różnym wieku, w różnych		dostarczanie bodźców, konтакты z ludźmi	Kontakty psa, Nauka zwi
13	stojach. Niech do mieszkania wchodzi Twój znajomy w		dostarczanie bodźców, konтакты z ludźmi, miejsca przebywania psa	Kontakty psa, Miejsca pr:
14	placzkach, kapturach, z parasolami, z łaskami. Jeżeli nie		dostarczanie bodźców, nowe sytuacje	Nauka zwierzaka, nowa :

Ilustracja 1.12. Zaznaczanie kilku wersów w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

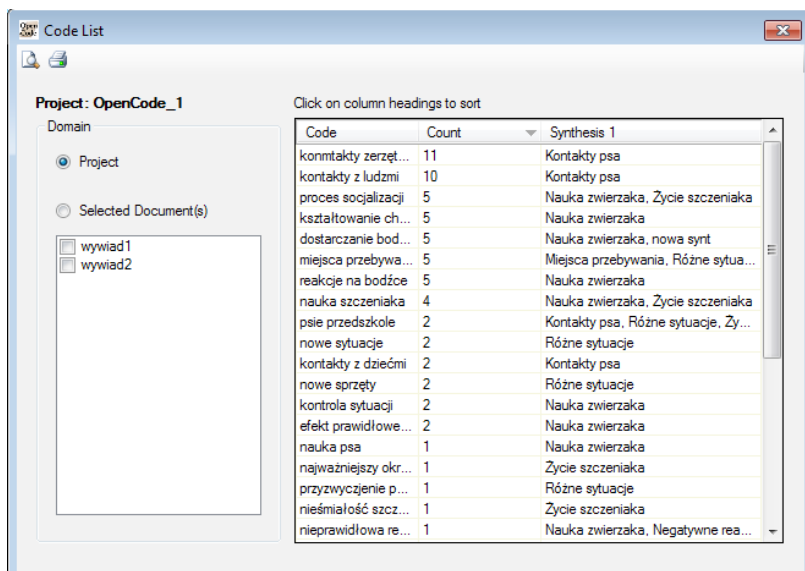
1.4.3. Przeglądanie listy kodów

Program OpenCode daje także możliwość przeglądania listy istniejących już kodów i związanych z nimi fragmentów dokumentów (*Code List*) oraz ich przeszukiwania (*Search By Code*). Obie funkcje są dostępne zarówno dla wybranego dokumentu tekstowego, jak również dla całej bazy danych.

Lista kodów zawiera kolumnę licznika, co oznacza liczbę wystąpień każdego kodu i pozwala na przegląd wszystkich kodów, jakie do tej pory utworzył badacz, oraz segmentów dokumentów, którym zostały one przypisane. W ten łatwy sposób można spojrzeć holistycznie na dotychczasowe wyniki pracy analitycznej i podejmować ewentualnie decyzję o wprowadzeniu określonych modyfikacji, polegających na przykład na zmianie nazwy bądź usunięciu kodu, jeśli jest to uzasadnione wynikami procesu interpretacyjnego danych bądź wynika po prostu z konieczności korekty pomyłki w nazwie kodu.

Lista kodów może być wyświetlana dla konkretnego bądź wszystkich dokumentów znajdujących się w projekcie.

Po wybraniu powyższej opcji, użytkownik może wybrać sposób wyświetlania: według częstotliwości (od najczęściej do najrzadziej występującego / bądź na odwrót) oraz alfabetycznie. Aby wybrać któryś z zadanych sposobów wyświetlania, należy kliknąć odpowiednią ilość razy na nagłówki kolumn, uzyskując w ten sposób określone zestawienie kodów.



Ilustracja 1.13. Użycie funkcji *Code List* w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

Uwaga: Kod przypisany raz w kilku kolejnych wierszach jest liczony jako jeden, pomimo że w oknie roboczym jest on widoczny w kilku liniach.

1.4.4. Przeszukiwanie kodów

Druga ze wspomnianych funkcji związanych z kodami to ich przeszukiwanie (*Search By Code*). Jak podkreślają twórcy oprogramowania, główną zaletą kodowania danych za pomocą oprogramowania komputerowego jest zdolność do wyszukiwania, która pozwala na znalezienie kodów, kategorii oraz segmentów tekstu z nimi związanych. Jest to o tyle ważne, że wyniki wyszukiwania mogą być kopiowane i wklejane do dokumentów programu Word, a funkcja przeszukiwania pozwala na łatwe wyszukiwanie cytatów w materiale, które można następnie umieścić w raporcie końcowym pracy.

Podobnie jak miało to miejsce przy funkcji wyświetlania listy kodów, także tutaj można wybrać, czy przeszukiwanie będzie dotyczyło pojedynczego dokumentu, czy całego projektu. W tym ostatnim przypadku poszczególne segmenty tekstów pochodzące z różnych dokumentów będą wyświetlane kolejno, co zdecydowanie ułatwia odleżenie konkretnych fragmentów tekstów w poszczególnych dokumentach.

Program OpenCode pozwala na dwa zasadnicze rodzaje przeszukiwania kodów. Pierwszy, prostszy polega na wyborze określonego pojedynczego kodu w polu listy wyboru kodów (*Code selection*). Można to zrobić za pomocą wyszukiwania dla jednego kodu lub kombinacji kodów. Wszystkie wystąpienia tego kodu pojawią się w oknie wyników.

Drugi sposób polega na wyborze z rozwijanej listy operatorów (OR, AND i AND NOT) i określeniu kombinacji kilku kodów służących do tworzenia formuły przeszukiwania. Używając wskazanych operatorów, możemy bardzo precyzyjnie określić sposób przeszukiwania, uzyskując wgląd w to, jakie segmenty danych zostały zakodowane kilkoma wskazanymi kodami jednocześnie (funkcja AND), po drugie, które kody współwystępują z kodem stanowiącym podstawę przeszukiwania (funkcja OR), a także wykluczyć te fragmenty tekstów, które zostały zakodowane określonym kodem (funkcja AND NOT).

Wskazane operatory można stosować jednocześnie, tworząc różne ich kombinacje i w ten sposób precyzyjnie określać parametry przeszukiwania kodów.

Aby rozpocząć nowe przeszukiwanie, należy użyć przycisku *Clear Selections* (Wyczyść zaznaczenia), co spowoduje usunięcie wcześniejszych ustawień parametrów przeszukiwania.

Wyniki przeszukiwania można następnie wydrukować (a wcześniej przejrzeć w podglądzie wydruku) oraz skopiować do schowka (kliknięcie w polu roboczym

Search By Code

Domain

Project

Selected document(s)

wywiad1

wywiad2

Selection

Clear Selection

OR

Select Code:

dobrze zachowanie w różnych sytuacjach

efekt prawidłowej socjalizacji

gabinet weterynaryjny

kontakty z rodziną

kontakty z dziećmi

kontakty z mężczyznami

kontrola sytuacji

kształtowanie charakteru

lekarz weterynaryjny

miejsca przebywania psa

najważniejszy okres szczeniaka

nauka psa

nauka szczeniaka

nazwa kodu

negatywne reakcje obcego

Code = dostarczanie bodźców OR kontakty z kobietami OR kontakty z ludźmi

LineNo	Text	Code
	Document: wywiad1	
2	właściwych kontaktów z ludźmi i innymi zwierzętami oraz	kontakty z rodziną, kontakty z ludźmi
6	się towarzystwem ludzi, może być zabierany w każde miejsce	efekt prawidłowej socjalizacji, kontakty z ludźmi
10	szczeniaka. To wtedy maluch powinien spotkać się z jak	dostarczanie bodźców, nauka szczeniaka
11	największą liczbą osób i innych psów. Jeżeli możesz	kontakty z rodziną, kontakty z ludźmi, nauka
12	zaangażuj spotkania z ludźmi w różnym wieku, w różnych	dostarczanie bodźców, kontakty z ludźmi
13	strojach. Niech do mieszkania wchodzi Twoi znajomi w	dostarczanie bodźców, kontakty z ludźmi, mi
14	plaszczach, kapturach, z parasolami, z laskami. Jeżeli nie	dostarczanie bodźców, nowe sytuacje
15	masz dzieci poproś o spotkanie z dziećmi sąsiadów,	dostarczanie bodźców, kontakty z dziećmi
16	znajomych. Psy często źle reagują na postawnych mężczyzn,	dostarczanie bodźców, kontakty z mężczyznami
17	szczególnie jeżeli domownicy to kobiety i dzieci. Dobieraj	kontakty z dziećmi, kontakty z kobietami
18	znajomych tak, by pokazać psu jak najlepiej różnych postaci.	kontakty z ludźmi, kształtowanie charakteru
20	dbać o to, by takie spotkania odbywały się w przyjaznej	kontakty z ludźmi, kontrola sytuacji
23	reakcjami na dzieci, ludzi starszych lub chorych.	kontakty z ludźmi, nieprawidłowa reakcja na
25	Równie ważne jak kontakty z ludźmi są kontakty z psami.	kontakty z rodziną, kontakty z ludźmi
30	właściciela czy Twój szczeniak może przywitać się z jego	kontakty z rodziną, kontakty z ludźmi
31	psiem. Spotykaj się jak najczęściej ze swoimi znajomymi	kontakty z ludźmi

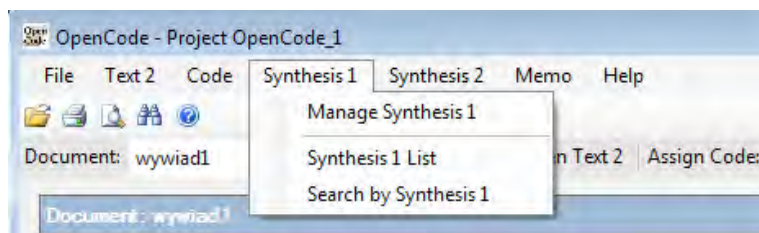
Ilustracja 1.14. Funkcja przeszukiwania kodów w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

w dowolnym miejscu na siatce przeszukiwania i naciśnięcie Ctrl + A spowoduje zaznaczenie wszystkich wierszy jednocześnie), używając standardowego zestawienia klawiszy klawiaturowych Ctrl + C, a następnie wkleić do edytora tekstu (np. MS Word).

1.4.5. Proces kategoryzowania – tworzenie i zarządzanie syntezami pierwszego stopnia

W programie OpenCode 4 mamy do dyspozycji funkcję syntetyzowania danych umożliwiającą wykonanie dwuetapowego klastrowania i konceptualizacji, niezależnie od przyjętego podejścia analitycznego. Służą do tego opcje *Synthesis 1* i *Synthesis 2* znajdujące się w menu głównym okna programu.



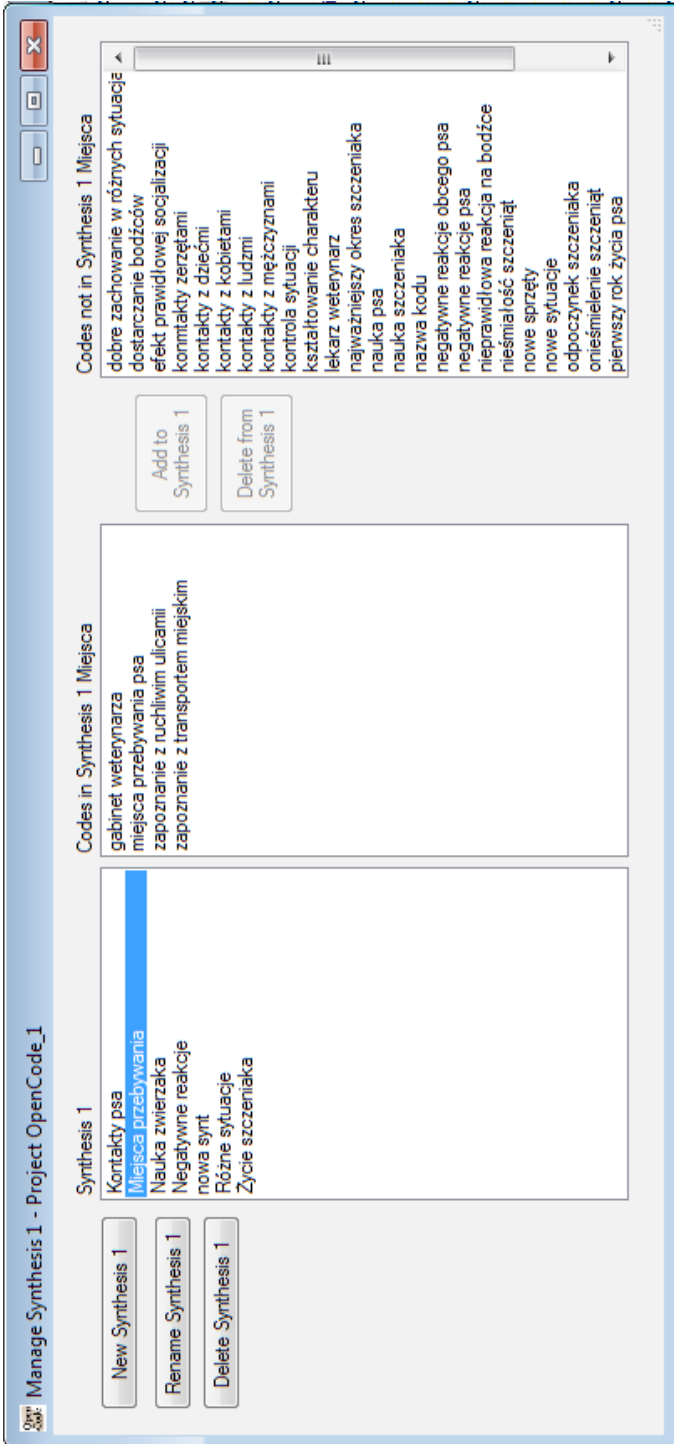
Ilustracja 1.15. Opcje syntetyzowania danych dostępne w menu głównym programu

Źródło: opracowanie własne na podstawie programu OpenCode 4

Kody utworzone przez badacza w procesie kodowania danych można pogrupować do konkretnych, zdefiniowanych przez niego kategorii. Operacje tego rodzaju wykonujemy, korzystając z funkcji *Synthesis 1* (Synteza 1) w menu głównym i wybierając jedną z trzech opcji: *Manage Synthesis 1* (Zarządzanie syntezą 1), *Synthesis 1 List* (Przeglądanie listy) oraz *Search By Synthesis 1* (Przeszukiwanie zawartości).

Pierwsza ze wskazanych funkcji *Manage Synthesis 1* (Zarządzanie syntezą 1) służy do tworzenia kategorii, przypisywania im określonych kodów, a także zmiany nazwy istniejącej już kategorii oraz ich ewentualnego usuwania. Po wywołaniu wspomnianej funkcji otwiera się okno z trzema kolumnami. Pierwsza kolumna zawiera nazwy *Synthesis 1*, w środkowej wyświetlane zostają kody przypisane do danej *Synthesis 1*, zaś trzecia kolumna zawiera wszystkie kody, jakie dotychczas zostały przez badacza wygenerowane w procesie analizy danych.

Utworzenie nowej *Synthesis 1* polega na naciśnięciu przycisku *New Synthesis 1* (Nowa synteza 1), a następnie wpisaniu jej nazwy. Przy czym należy pamiętać, że tak jak ma to miejsce w przypadku nazw kodów, także nazwy Syntezy nie mogą być dłuższe niż 60 znaków. Nie ma natomiast ograniczeń co do stosowania polskich liter oraz innych znaków, za wyjątkiem przecinków.



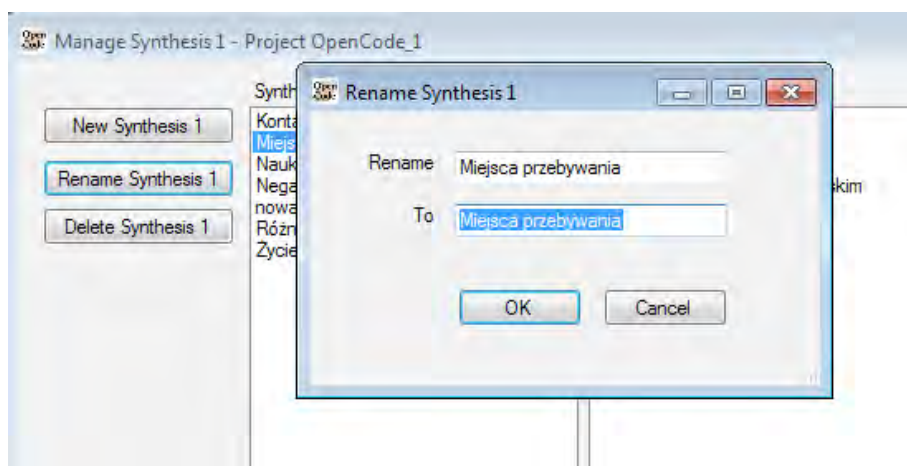
Ilustracja 1.16. Przypisywanie kodów do syntezy (kategorii) w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

Kolejną czynnością jest przypisanie wybranych przez badacza kodów do utworzonej Syntezy 1. W tym celu należy zaznaczyć (podświetlić) dany kod znajdujący się w trzeciej kolumnie okna dialogowego, a następnie użyć przycisku „<”. Aby usunąć kod z danej Syntezy, trzeba wybrać ów kod z listy kodów przypisanych do Syntezy i nacisnąć przycisk „>”. Zarówno w pierwszym, jak i drugim przypadku można wybrać więcej niż jeden kod, naciskając i przytrzymując klawisz Ctrl.

Uwaga: Dany kod może być przypisany do wielu kategorii (syntez) jednocześnie.

Ponieważ analiza i interpretacja danych jest procesem, w trakcie którego badacz może wielokrotnie modyfikować dotychczasowe ustalenia, program OpenCode oferuje pomocne w tym zakresie opcje zmiany nazw istniejącym już kategoriom (*Rename Synthesis 1*) jak również ich usuwania (*Delete Synthesis 1*). Niemniej istotna jest w tym zakresie możliwość wyłączenia lub włączenia nowych kodów do danej kategorii.



Ilustracja 1.17. Zmiana nazwy kategorii w programie OpenCode 4

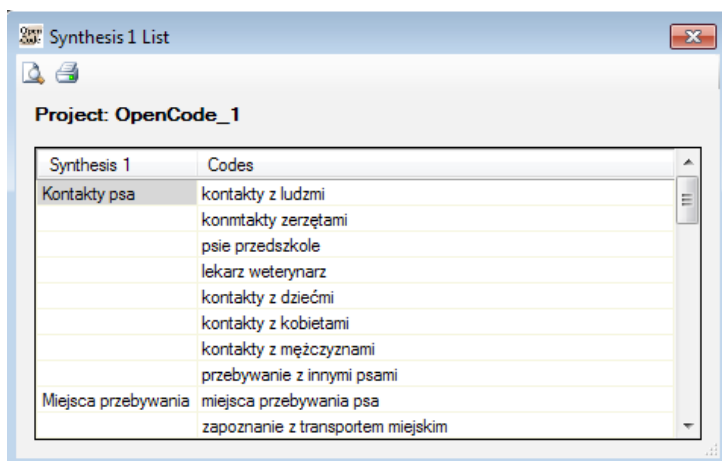
Źródło: opracowanie własne na podstawie programu OpenCode 4

1.4.6. Przeszukiwanie syntez pod względem przypisanych im kodów

Dużym ułatwieniem w procesie analitycznym są dwie kolejne funkcje, jakie można wykonywać, operując kategoriami. Są to wspomniane wcześniej: przeglądanie listy syntez (*Synthesis 1 List*) oraz przeszukiwania zawartości syntez (*Search By Synthesis 1*).

Badacz, który chce zapoznać się z listą wszystkich dotychczas utworzonych syntez oraz z przypisanymi do nich kodami, powinien skorzystać z pierwszej ze wspomnianych funkcji. W wyświetlonym oknie po lewej stronie znajdują się nazwy syntez, po prawej zaś przypisane im kody. Kolejność wyświetlania syntez można zmienić, naciskając na główkę tabeli.

Badacz, który chce dokładniejszych informacji, dotyczących zestawienia syntez, przypisanych do nich kodów oraz listy wierszy, do których zakodowania zostały użyte poszczególne kody, może użyć funkcji przeszukiwania zawartości syntezy (*Search By Synthesis 1*). W otwartym oknie, po lewej stronie mamy wykaz istniejących syntez, po prawej zaś okno z podziałem na: tekst, numer linii, kody, syntezy. W ten sposób badacz uzyskuje dostęp do informacji o treści zakodowanych segmentów tekstu z podziałem na poszczególne dokumenty, numerach linii, w których znajdują się te fragmenty, a także przypisanych im kodach oraz zestawieniu syntez, do których należą wyświetlone kody (pamiętajmy bowiem, że dany kod może być przypisany do więcej niż jednej kategorii).

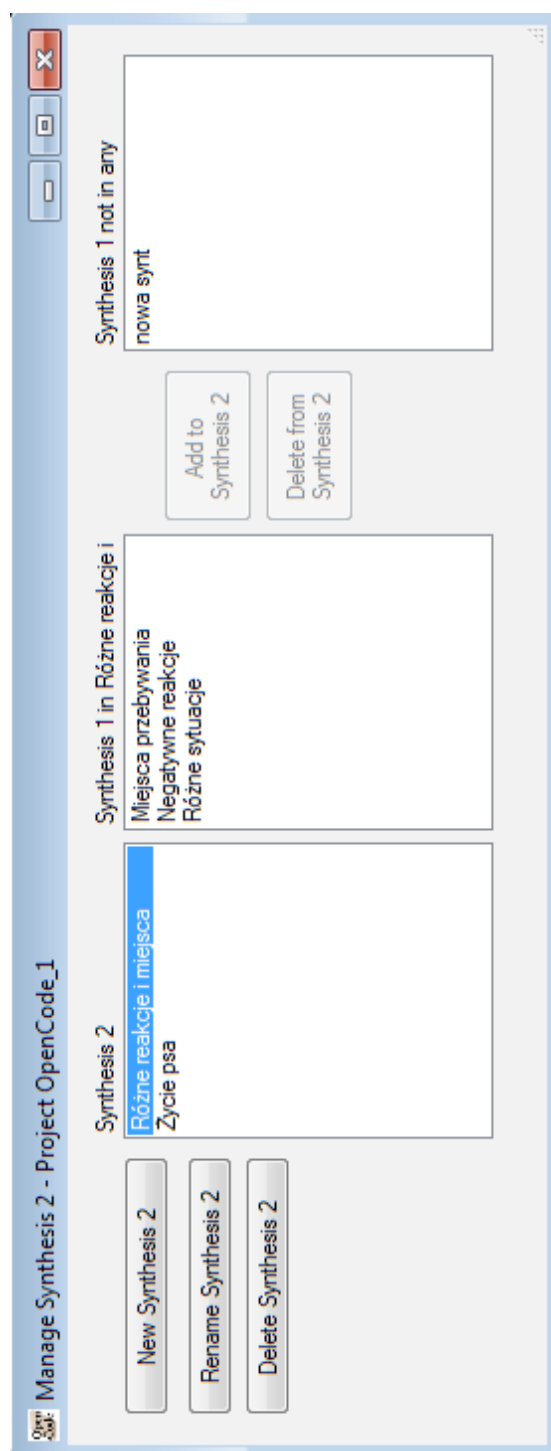


Ilustracja 1.18. Przeglądanie listy kodów za pomocą *Search By Synthesis 1*

Źródło: opracowanie własne na podstawie programu OpenCode 4

1.4.7. Tworzenie i zarządzanie syntezą drugiego stopnia

W programie oprócz syntez pierwszego stopnia można także tworzyć bardziej rozbudowane formuły, opierające się na interpretacji danych, przechodząc na kolejny poziom uogólniania wyników. Służą do tego syntezy drugiego stopnia, które pozwalają na grupowanie syntez stopnia pierwszego. Zasadniczo nie ma jakiejś strukturalnej różnicy pomiędzy syntezą 1 i syntezą 2. Nie jest jedynie możliwe



Ilustracja 1.20. Okno *Synthesis 2* w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

uwzględnienie tych samych pojęć syntezy pierwszego stopnia w kilku koncepcjach syntezy stopnia wyższego.

Utworzenie nowej *Synthesis 2* jest tożsame z czynnościami wykonywanymi podczas tworzenia *Synthesis 1* i polega na naciśnięciu przycisku *New Synthesis 2* (Nowa Synteza 2), a następnie wpisaniu jej nazwy. Podobnie jest z możliwością zmiany nazw istniejącym już syntezom (*Rename Synthesis 2*) jak również ich usuwania (*Delete Synthesis 2*).

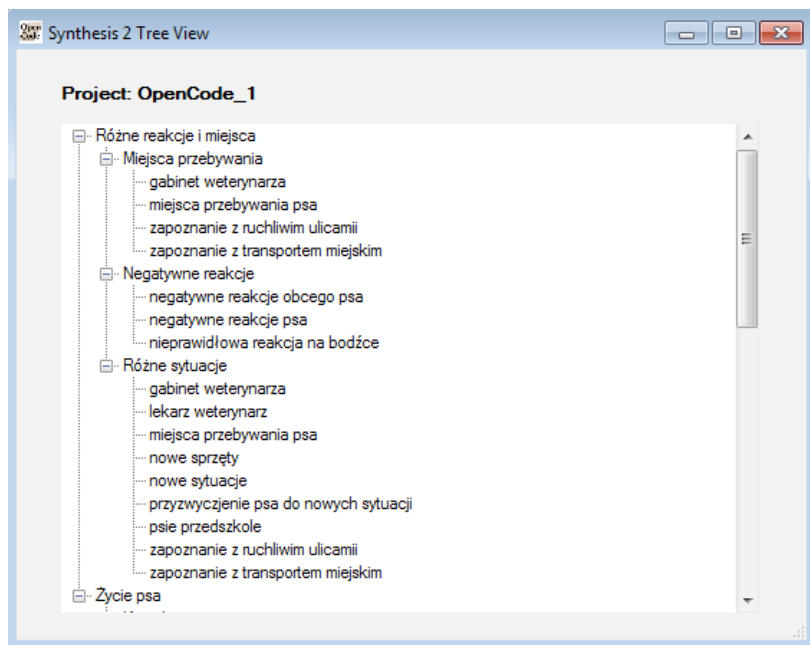
Również przeglądanie listy syntez (*Synthesis 2 List*) jest tożsame z tym, co badacz może wykonać w trakcie pracy nad syntezami pierwszego stopnia. Przy czym w tym wypadku wyświetlają się nam informacje zarówno o kodach przypisanych do Syntezy 1, jak i Syntezy pierwszego stopnia przypisane do Syntezy 2.

Tym, co różni opcje dostępne w ramach funkcji *Synthesis 1* i *Synthesis 2*, jest zestawienie kodów, syntez pierwszego i drugiego stopnia, w postaci widoku drzewa (*Synthesis 2 Tree View*) dostępne w menu *Synthesis 2*. Dzięki tej opcji badacz ma możliwość obejrzenia rezultatów swojej pracy analitycznej w bardziej przystępny i holistyczny sposób.

Synthesis 2	Synthesis 1	Code
Różne reakcje i miejsca	Miejsca przebywania	gabinet weterynarza
Różne reakcje i miejsca	Miejsca przebywania	miejsca przebywania psa
Różne reakcje i miejsca	Miejsca przebywania	zapoznanie z ruchliwymi ulicami
Różne reakcje i miejsca	Miejsca przebywania	zapoznanie z transportem miejskim
Negatywne reakcje	Negatywne reakcje	negatywne reakcje obcego psa
Negatywne reakcje	Negatywne reakcje	negatywne reakcje psa
Negatywne reakcje	Negatywne reakcje	nieprawidłowa reakcja na bodźce
Różne sytuacje	Różne sytuacje	gabinet weterynarza
Różne sytuacje	Różne sytuacje	lekarz weterynarz
Różne sytuacje	Różne sytuacje	miejsca przebywania psa
Różne sytuacje	Różne sytuacje	nowe sprzęty
Różne sytuacje	Różne sytuacje	nowe sytuacje
Różne sytuacje	Różne sytuacje	przywroczenie psa do nowych sytuacji

Ilustracja 1.21. Okno *Synthesis 2 List*

Źródło: opracowanie własne na podstawie programu OpenCode 4



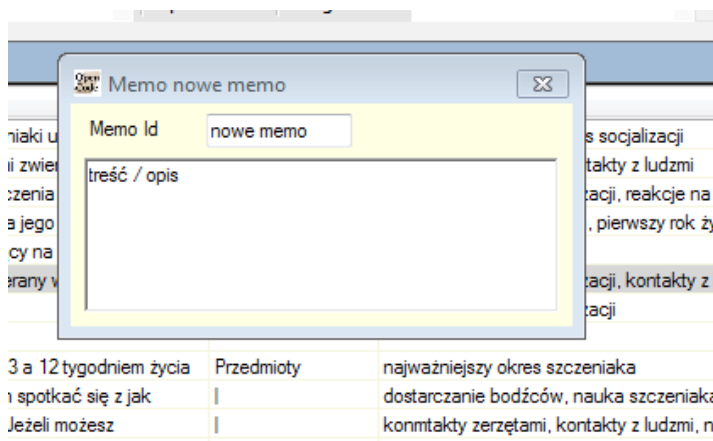
Ilustracja 1.22. Widok drzewa (*Tree View*) zbudowanego z kodów oraz syntez pierwszego i drugiego stopnia

Źródło: opracowanie własne na podstawie programu OpenCode 4

1.4.8. Tworzenie memo

Badacz, który podczas generowania kodów i kategorii będzie chciał zapisać swoje pomysły analityczne dotyczące powiązań między kategoriami i tworzonych na tej podstawie hipotez bądź po prostu krótkich informacji na temat kodów i kategorii, może skorzystać z funkcji tworzenia memo. Notatki w formie memo można utworzyć, przypisując je do konkretnych segmentów tekstu, a także edytować i modyfikować na dwa sposoby. Można to uczynić, korzystając z menu głównego (*Memos*) bądź z menu kontekstowego wywoływanego prawym przyciskiem myszy po zaznaczeniu określonego fragmentu tekstu i wybraniu opcji *New Memo to selected lines* (Utwórz nowe memo w zaznaczonych wersach). W tym ostatnim przypadku kursor musi się znajdować w obszarze kolumny memo. W otwartym oknie należy wpisać nazwę memo oraz treść notatki, którą w każdej chwili badacz może edytować, korzystając z opcji *Show (Edit) Memo* (Pokaż/Edytuj memo). Po stworzeniu notatki należy zamknąć okno dialogowe bądź pozostawić je nadal otwarte. Wszystkie wpisane informacje zostaną automatycznie zapisane w bazie danych projektu.

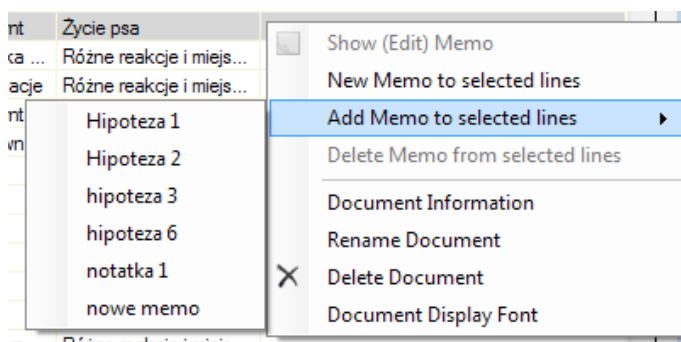
Uwaga: W programie OpenCode nie ma funkcji *Save* (Zapamiętaj). Wynika to z tego, że program został tak pomyślany, aby wszelkie wprowadzane przez użytkownika dane oraz zmiany były zapisywane automatycznie.



Ilustracja 1.23. Tworzenie memo w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

W programie OpenCode można po zaznaczeniu wybranych linii wyświetlać istniejące memo, dokonać ich modyfikacji bądź je usuwać. Wspomniane opcje można wywołać zarówno z menu głównego *Memo*, jak również menu kontekstowego. Dodatkową opcją dostępną z poziomu menu kontekstowego jest przypisanie notatki w formie memo określonego segmentowi tekstu, spośród wybranych memo z rozwijanej listy (*Add Memo to selected lines*).



Ilustracja 1.24. Przypisywanie memo z rozwijanej listy w menu kontekstowym w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

1.5. Wyszukiwanie danych w tekście

OpenCode należy do rodziny programów „kodu i wyszukiwania” i z tego względu posiada dość rozbudowane opcje przeszukiwania. O dwóch z nich wspomniano już wcześniej. Były to: opcja przeszukiwania kodów oraz zawartości kategorii. Oprócz tego OpenCode pozwala na wyszukiwanie słów oraz ich fragmentów w treści dokumentów projektu.

Funkcja przeszukiwania dokumentów jest podobna do tej, z jaką możemy się spotkać w programie MS Word i pozwala na szybkie odnalezienie słów oraz fraz w aktualnie otwartym pliku dokumentu. Aby użyć tej funkcji przeszukiwania, należy wybrać ikonkę lornetki na pasku narzędzi, skorzystać z menu głównego *File/Find in Document Text* (Plik/Odszukaj w dokumencie) bądź skrótu klawiszowego Ctrl + F.

Panel opcji przeszukiwania otwiera się w formie pomarańczowego paska w górnej części okna roboczego. Tekst dokumentu może być przeszukiwany linia po linii (*Find after line*) lub w całości (*Find All*). Linie, które zawierają wyszukiwane słowo bądź frazę, zostają podświetlone na szaro. O liczbie wierszy, w których znajduje się dane słowo lub fraza, użytkownik może się dowiedzieć z informacji zamieszczonej w lewym dolnym rogu okna programu.

Wyłączenie opcji wyszukiwania następuje poprzez ponowne naciśnięcie ikonki lornetki na pasku narzędzi bądź skorzystanie z menu głównego *File/Find in Document Text*.

1.6. Funkcje przeglądania, eksportowania i drukowania danych

Program OpenCode pozwala także na bieżącą kontrolę dotychczasowego stanu zaawansowania analizy prowadzonej przez badacza. Bardzo pomoce w tym zakresie jest operowanie opcją podglądu wydruku (*Print Preview*), którą można wywołać, korzystając z ikonki znajdującej się na pasku narzędzi bądź w menu głównym *File*. Ponadto ikona podglądu wydruku i towarzysząca jej ikonka opcji drukowania znajdują się w prawym górnym rogu okienek przeszukiwania i przeglądania listy kodów oraz kategorii. Przy czym w tych przypadkach użycie ikony powoduje automatyczne wyświetlenie się podglądu wydruku aktualnie wykonywanych czynności. Jeśli natomiast skorzystamy z opcji podglądu wydruku za pośrednictwem ikony znajdującej się na pasku narzędzi bądź wybierzemy ją z menu głównego (*File/Print Preview*), wówczas otrzymamy możliwość ustalenia, jakie informacje mają być w ten sposób wyświetlane.

W otwartym oknie wyboru opcji podglądu wydruku badacz może określić, jakie dane mają być prezentowane w ramach danego dokumentu, jakie będą tworzyły zestawienie oraz jaki będzie format tekstu, a dokładnej odstępów między wierszami.

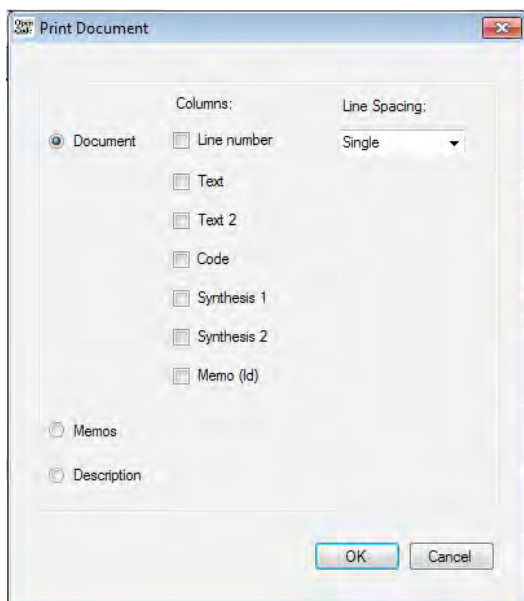
OpenCode - Project OpenCode_1					
File	Text 2	Code	Synthesis 1	Synthesis 2	Memo
Document: wywiad1	Open Text 2	Assign Code:	Add	Assigned codes: konmtakty zerzetami	Remove
Document: wywiad1					
soc	Find after line 65	Find All			
Text	Text 2	Code	Synthesis 1	Synthesis 2	Memo
19	Zawsze musisz kontrolować sytuację. Należy bezwzględnie	kontrola sytuacji	Nauka zwierzaka	Zycie psa	
20	dbać o to, by jakie spotkania odbywały się w przyjaznej	kontakty z ludźmi, kontrola sytuacji	Kontakty psa, Nauka zwierzaka	Zycie psa	
21	atmosferze i były mile dla psaka. Każde złe doświadczenie	nieprawidłowa reakcja na bodźce	Nauka zwierzaka, Negatywne reakcje	Różne reakcje i meje...	
22	z tego okresu odbija się w przyszłości nieprawidłowymi	nieprawidłowa reakcja na bodźce	Nauka zwierzaka, Negatywne reakcje	Różne reakcje i meje...	
23	reakcjami na dzieci, ludzi starszych lub chorych.	kontakty z ludźmi, nieprawidłowa rea...	Kontakty psa, Nauka zwierzaka, Negatyw...	Różne reakcje i meje...	
24					
25	Równie ważne jak kontakty z ludźmi są kontakty z psami.	konmtakty zerzetami, kontakty z ludźmi	Kontakty psa	Zycie psa	
26	Dobre jest przedstawić szczeniaka psom z sąsiedztwa.	nauka szczeniaka, przebywanie z inn...	Kontakty psa, Nauka zwierzaka, Zycie szcz...	Zycie psa	
27	Przedstawiając szczeniaka psom dorosłym musisz zachować	kontrola sytuacji, przebywanie z inny...	Kontakty psa, Nauka zwierzaka, Zycie szcz...	Zycie psa	
28	większą ostrożność. Pamiętaj, że nie każdy dorosły pies	konmtakty zerzetami, kontrola sytuacji...	Kontakty psa, Nauka zwierzaka, Negatyw...	Różne reakcje i meje...	
29	jest ciekawym w stosunku do szczeniąt. Zawsze zapytaj	konmtakty zerzetami, negatywne reak...	Kontakty psa, Negatywne reakcje	Różne reakcje i meje...	
30	właściciela czy Twój szczeniak może przywitać się z jego	konmtakty zerzetami, kontakty z ludźmi	Kontakty psa	Zycie psa	
31	psiem. Spotykaj się jak najczęściej ze swoimi znajomymi	kontakty z ludźmi	Kontakty psa	Zycie psa	
32	posadającymi psy. W parku rozejrzyj się, czy ktoś nie	kontakty z ludźmi, miejsca przebywani...	Kontakty psa, Miejsca przebywania, Różne ...	Różne reakcje i meje...	
33	spacując z psakiem w podobnym wieku - może to być	konmtakty zerzetami	Kontakty psa	Zycie psa	
34	wyjemity kompan. Twojego szczeniaka. Pomocne w	konmtakty zerzetami	Kontakty psa	Zycie psa	
35	socializacji z innymi psami są tzw. przedszkola dla psów".	konmtakty zerzetami, psie przedszkole	Kontakty psa, Różne sytuacje, Zycie szczen...	Różne reakcje i meje...	
36	Psaki uczą się podstawowych poleceń, a przede wszystkim	kształtowanie charakteru	Nauka zwierzaka	Zycie psa	
37	poznają się i uczą wzajemnych kontaktów. Pamiętaj jednak,	konmtakty zerzetami, kształtowanie c...	Kontakty psa, Nauka zwierzaka	Zycie psa	
38	że nawet najlepsza szkółka nie wystarczy - większość pracy	proces socializacji, psie przedszkole	Kontakty psa, Nauka zwierzaka, Różne sytu...	Różne reakcje i meje...	

Ilustracja 1.25. Wyszukiwanie słów w dokumencie projektu w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

Ta ostatnia opcja została udostępniona, aby po wydrukowaniu badacz, przeglądając i analizując dane, mógł w wygodny sposób nanosić swoje uwagi i zapisywać dodatkowe informacje bezpośrednio na wersji papierowej pomiędzy wersjami.

W pierwszej kolejności badacz musi dokonać wyboru pomiędzy trzema głównymi opcjami podglądu: całego dokumentu z wyszczególnieniem określonych zestawień, zawartości memo czy informacji o dokumencie. Następnie, o ile badacz wybrał pierwszą z wymienionych opcji, musi wskazać, jakie dokładnie dane mają się znaleźć w poszczególnych kolumnach w podglądzie wydruku (numery linii, tekst dokumentu, nazwy memo, kody, kategorie). Cechą charakterystyczną jest to, że rodzaj wyświetlanych w ten sposób danych jest dowolny. Ostatnia kwestia do ustalenia to odstęp między wersami (pojedyncze, półtora wersu bądź podwójne odstęp).

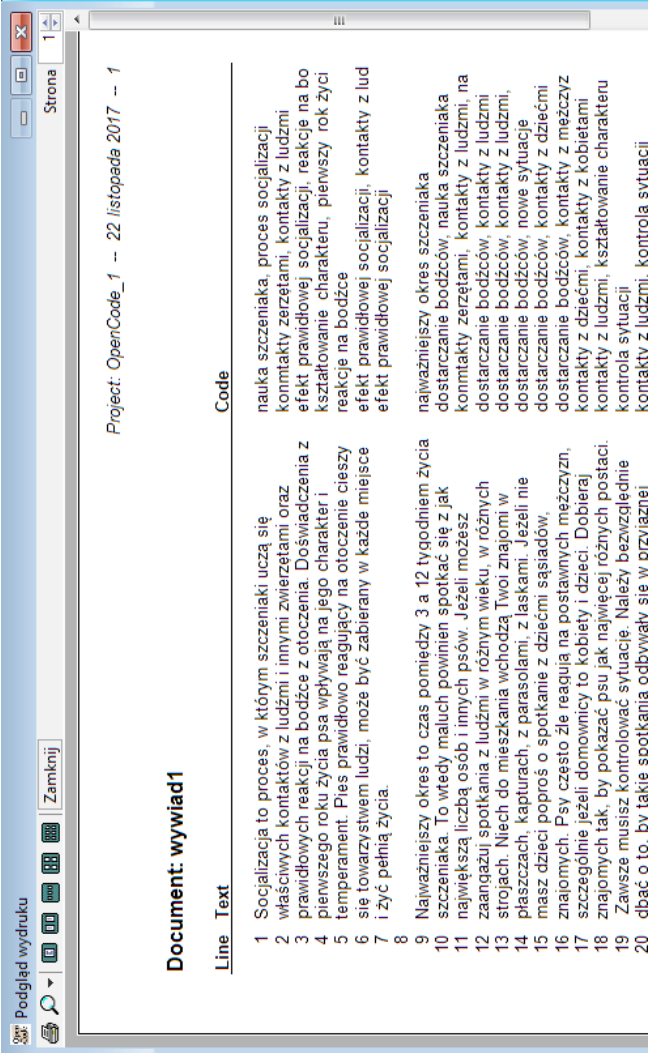


Ilustracja 1.26. Opcje podglądu wydruku w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

Badacz dzięki opcji podglądu wydruku może nie tylko sprawdzić określony dokument przed jego wydrukiem, ale także w przejrzysty i wygodny sposób zapoznać się z wybranymi zestawieniami dotyczącymi kodów, syntez, memo, zaimportowanych materiałów tekstowych czy informacji opisujących dany dokument.

W otwartym oknie podglądu wydruku widoczne są, poza uprzednio wybranymi, takie informacje, jak nazwa projektu, data utworzenia danej wersji oraz numeracja stron. Wszystkie wymienione informacje znajdują się w górnym prawym rogu okna, przy czym nazwa projektu oraz data utworzenia wersji do wydruku

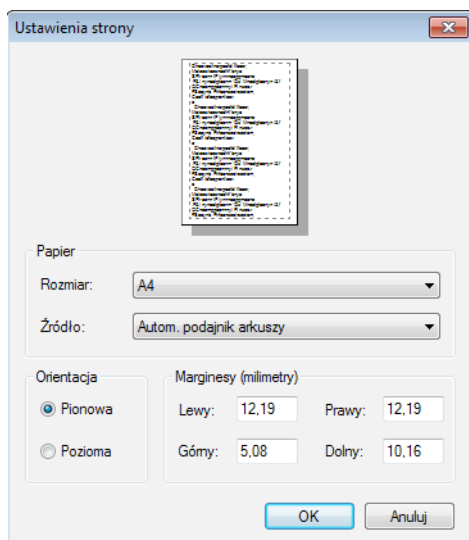


Ilustracja 1.27. Podgląd wydruku dokumentu w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

znajduje się wyłącznie na pierwszej stronie. Przechodzenie pomiędzy stronami umożliwia zaś znajdująca się w najwyższym punkcie prawego górnego rogu okna opcja poruszania się pomiędzy stronami przygotowanymi do wydruku.

Oprócz funkcji podglądu wydruku jest także opcja *Print Settings...* (Ustawienia wydruku) dostępna w menu *File*, dzięki której użytkownik programu może zoptymalizować ustawienia drukowania. W wyświetlonym oknie można ustawić między innymi rozmiar marginesów, rodzaj rozmiaru strony czy jej orientację (pionowa/pozioma). Jest to szczególnie przydatne, gdy ilość danych, jakie badacz chce umieścić w wersji drukowanej, jest na tyle obszerna, że ustawienie standardowe marginesów czy pionowej orientacji strony uniemożliwia wyświetlenie wszystkich niezbędnych informacji.



Ilustracja 1.28. Opcje ustawień drukowania w programie OpenCode 4

Źródło: opracowanie własne na podstawie programu OpenCode 4

Opisane powyżej funkcje stanowią niezwykle pomocne narzędzia w procesie prowadzenia analizy i interpretacji danych tekstowych przede wszystkim zgodnie z metodologią teorii ugruntowanej. Jednak ze względu na ograniczone możliwości techniczne oprogramowania dalsze postępowanie analityczne wymaga powrotu do tradycyjnych form pracy badawczej. Na szczęście program OpenCode oferuje bardzo dogodne opcje podglądu i wydruku poszczególnych elementów projektu (kodów, kategorii, memo) oraz ich zestawień. Dzięki temu badacz zyskuje możliwość kontynuowania analizy (np. tworzenia map i diagramów integracyjnych oraz teoretycznego opisu) na podstawie wcześniej zakodowanych i posegregowanych w programie danych.

1.7. Uwagi końcowe

Poza zaprezentowanymi zaletami programu jako narzędzia wspomagającego analizę danych jakościowych można wskazać na kilka ograniczeń, które mogą niekorzystnie wpływać na pracę analityczną.

Przed wszystkim wydaje się, że sposób prezentowania (wyświetlania) danych narzuca określony i właściwie niepodlegający modyfikacjom układ treści materiałów tekstowych, automatycznie dzieląc zawartość dokumentu na wersy składające się zawsze z równej liczby sześćdziesięciu znaków.

Ponadto narzucony podział tekstów może być w niektórych przypadkach dość sztuczny, co uwydatnia się szczególnie w momencie kodowania danych. W konsekwencji określone porcje (fragmenty) dokumentu, którym badacz nadaje kody, mogą rozpoczynać się i kończyć w różnych momentach poszczególnych wersów, a jeśli badacz zakoduje różnymi kodami dwa fragmenty znajdujące się w tym samym wersie, to taka konstrukcja staje się miejscami mało czytelna.

Program OpenCode nie daje także możliwości dopasowania sposobu dostępu do poszczególnych funkcji zgodnie z potrzebami użytkownika (na przykład zmiany bądź poszerzenia gamy funkcji otwieranych z poziomu ikon na pasku narzędzi), co oznacza, że także układ interfejsu jest tylko w niewielkim stopniu modyfikowalny (można na przykład zmienić szerokość kolumn).

Znacznym ograniczeniem programu wydaje się być możliwość importowania dokumentów jedynie w formacie .txt, co uniemożliwia pracę na materiale wcześniej sformatowanym w innym programie służącym do edycji tekstu.

Co więcej, program OpenCode nie pozwala na edytowanie tekstu już zaimportowanego². W tekście nie można wprowadzać żadnych zmian, a więc wszelkie błędy, które badacz odkryje po zaimportowaniu dokumentu, nie podlegają korekcie. Jest ona możliwa jedynie w tekście jeszcze niezaimportowanym, a zatem badacz musi po umieszczeniu dokumentu w bazie danych bardzo dokładnie sprawdzić, czy tekst jest poprawny i nie zawiera usterek. Jeśli takie się pojawiają, wówczas pozostaje jedynie usunąć plik i ponownie go zaimportować, już po wprowadzeniu stosownej korekty.

Ograniczeniem z punktu widzenia analizy danych jest brak możliwości wprowadzania hierarchii kategorii z wyszczególnieniem kategorii nadrzędnych i podrzędnych (podkategorii), a także ich własności (stąd o stopniu „ogólności” samych kategorii musi pamiętać badacz, ponieważ nie można oznaczyć ich

² Istnieje wprawdzie możliwość modyfikacji tekstu, która ogranicza się jednak do zmian wielkości, kroju oraz rodzaju czcionki, jednak zmiany te dotyczą całego tekstu, a nie jego fragmentu, stąd ich niewielka użyteczność w procesie analizy (np. nie ma możliwości wyróżnienia porcji danych). Służą one raczej zwiększeniu komfortu pracy użytkownika (np. poprzez zwiększenie rozmiaru/wielkości czcionki dokumentu).

miejsca w hierarchii ani określić znaczenia, jakie analityk nadaje poszczególnym kategoriom, np. zgodnie z paradygmatem kodowania metodologii teorii ugruntowanej).

Pewnym ograniczeniem jest również brak możliwości prowadzenia dalszych etapów analizy, takich jak łączenie kategorii, poszukiwanie związków między nimi czy określanie charakteru takiej zależności (tutaj pewną pomocą służą mema, gdzie badacz może zapisać informacje o związkach między kategoriami, a także o hipotezach, jakie wyłaniają się podczas interpretacji danych).

Pomimo wspomnianych ograniczeń, jakie wiążą się z użytkowaniem programu OpenCode, należy jednak pamiętać, że jest to narzędzie rodzaju *code-and-retrieve* (Bieliński i in. 2007: 93–94) i jako takie spełnia swoje zadanie w zupełności. Jeżeli zaś oczekujemy dodatkowych i bardziej rozbudowanych funkcji ze strony oprogramowania, na przykład przy tworzeniu związków i zależności między kategoriami czy budowaniu graficznych prezentacji analitycznych dokonań badacza (w postaci grafów bądź diagramów), powinniśmy poszukać programu bardziej zaawansowanego technologicznie i metodologicznie (przykładem takich programów są między innymi Atlas.ti, NVivo czy Maxqda). Jednak gdy zależy nam przede wszystkim na zapanowaniu nad obszerną bazą materiałów, a przy tym zamierzamy korzystać z danych tekstowych, to program OpenCode jest bardzo dobrą alternatywą dla odpłatnego oprogramowania CAQDAS, którego rozbudowane funkcje nie zawsze są nam jako badaczom potrzebne w realizacji konkretnego projektu badawczego.

OpenCode to oprogramowanie stworzone na potrzeby metodologii teorii ugruntowanej, do analizy tekstowych danych jakościowych, umożliwiające różnorodne kodowanie, zestawianie i porównywanie kategorii analitycznych, pisanie memo i w tym zakresie godne polecenia dla potencjalnych użytkowników. Zwłaszcza jeśli dopiero rozpoczynamy swoją przygodę z oprogramowaniem wspomagającym analizę danych jakościowych.

Warto nadmienić, że językiem programu jest angielski, ale dane (tekst) i kodowanie może być prowadzone w dowolnym języku. Stąd program może być z powodzeniem wykorzystywany przez badaczy posługujących się rodzimym językiem, którzy zmierzają opracowywać dane w ich oryginalnej wersji.

Program OpenCode 4 został stworzony i udostępniony przez ICT Services and System Development and Division of Epidemiology and Global Health, University of Umeå, Sweden.

Zarówno sam program, jak i dodatkowe informacje o nim można pobrać ze strony: <http://www.phmed.umu.se/english/divisions/epidemiology/research/open-code/?languageId=1>

Rozmiar pliku 0,81 MB (wersja OpenCode 4).

Przed instalacją programu należy upewnić się, czy na komputerze posiadamy oprogramowanie .Net Framework. Jeśli nie, można je pobrać z tej samej strony bądź innego źródła internetowego.

Systemy, pod którymi można uruchomić program, to Windows XP i nowsze (w wersji 32 i 64-bitowej)³.

Przy opracowywaniu materiału o programie posłużono się instrukcją użytkownika Klas Göran Sahlen i zespół z Department of Nursing and Department of Public Health and Clinical Medicine Umeå University, Szwecja.

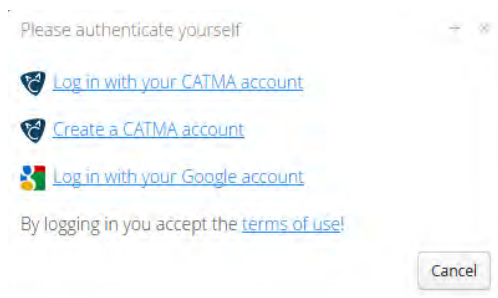
³ Zgodnie z informacjami autorów programu.

2. CATMA – profesjonalne narzędzie do analizy tekstu

CATMA (Computer Assisted Textual Markup and Analysis) jest praktycznym i intuicyjnym narzędziem dla badaczy zajmujących się analizą tekstów. Użytkownicy CATMA mogą połączyć hermeneutyczne podejście do analizy tekstu z możliwościami środowiska cyfrowego oferowanego przez oprogramowanie. Od wersji czwartej CATMA jest implementowany jako aplikacja internetowa (w przeciwieństwie do wcześniejszych generacji programu). Wśród najważniejszych cech oprogramowania wymienić można: obsługę cyfrowego tekstu w prawie każdym języku; pełną integrację funkcji do opisu i analizy danych w przeglądarce internetowej; współpracę za pośrednictwem sieci internetowej umożliwiającą łatwą wymianę dokumentów, adnotacji i znaczników; dowolnie definiowalne lub predefiniowane znaczniki, które można również udostępniać; interaktywne wyszukiwanie w języku naturalnym poprzez tekst, korpus i adnotacje; automatyczne, statystyczne i nie-statystyczne funkcje analityczne; wbudowaną wizualizację wyników wyszukiwania i analiz; analizę skomplikowanych korpusów tekstowych w jednym kroku. Program daje też możliwość pracy pojedynczego badacza lub współpracy w czasie rzeczywistym w ramach zespołu.

2.1. Rozpoczynanie pracy w środowisku programu CATMA

Pierwszą czynnością niezbędną do pracy w programie jest zalogowanie się.



Ilustracja 2.1. Panel logowania/utworzenia konta w programie CATMA

Źródło: opracowanie własne na podstawie programu CATMA

W tym celu należy odwiedzić stronę <http://portal.catma.de/catma/>. Czynność logowania wykonuje się za pośrednictwem posiadanego konta e-mail (Google-Mail). W tym wypadku wystarczy wpisać swój adres e-mail i hasło. Jeżeli jednak nie mamy konta mailowego, wówczas należy skorzystać z opcji bezpośredniego utworzenia konta użytkownika w ramach programu CATMA.

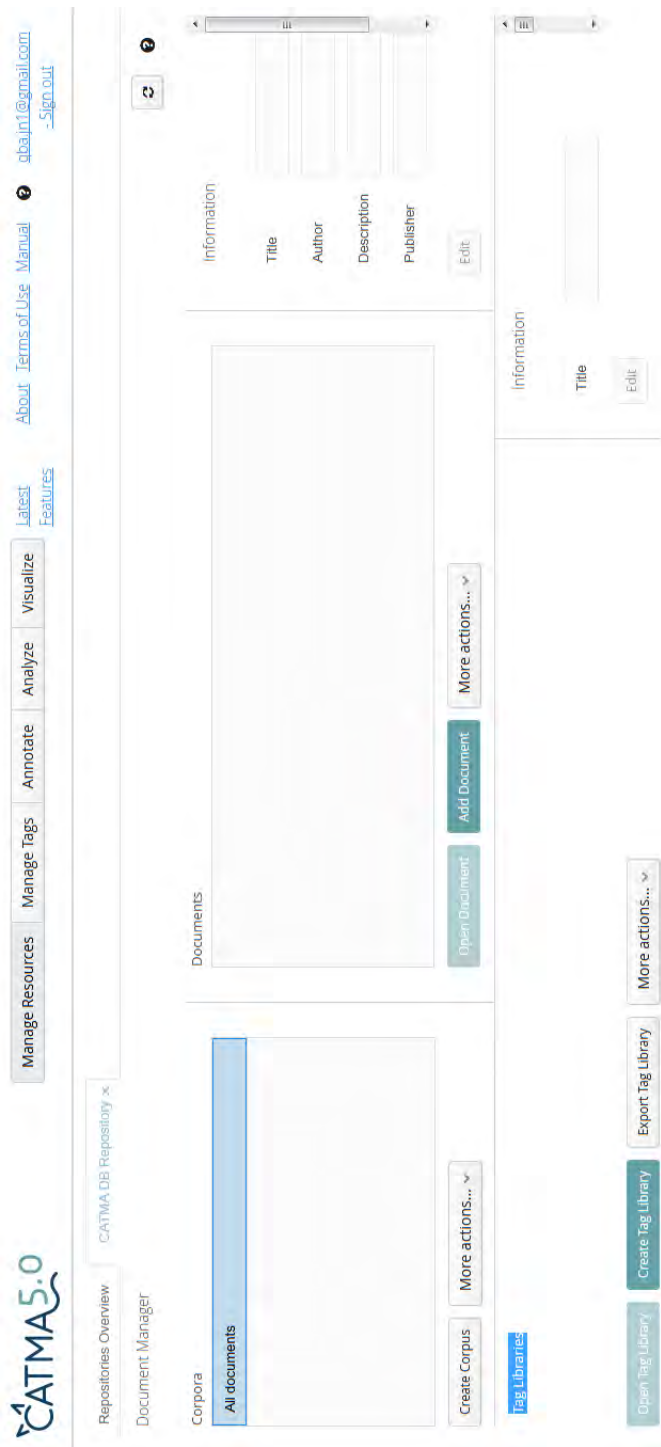
2.1.1. Dodawanie tekstów

Po zalogowaniu się do programu naszym oczom ukazuje się okno Menedżera programu, które podzielone jest na kilka kluczowych obszarów.

W lewym górnym rogu znajduje się lista korpusów (*Corpora*). Sekcja w środkowej górnej części okna pokazuje listę dokumentów (*Documents*). W prawym górnym rogu Menedżera repozytorium znajdują się informacje o poszczególnych dokumentach lub kolekcjach znaczników (m.in. tytuły, autorzy itp.). W lewej dolnej części okna wyświetla jest Biblioteka znaczników (*Tag Libraries*). Natomiast w dolnej prawej części Menedżera repozytorium widnieją informacje na temat tagów (*Tagsets*).

Pracę w programie rozpoczynamy od dodania dokumentów, co wykonujemy, klikając na przycisk *Add Document* znajdujący się w centralnym punkcie okna Menedżera. Program pozwala na dodanie zarówno pliku dostępnego za pośrednictwem Internetu poprzez podanie odpowiedniego adresu URL lub przesłanie pliku bezpośrednio z komputera. Jeśli wybrana zostanie druga opcja, wówczas należy kliknąć na przycisk *Upload local file* (Prześlij lokalny plik), a po zakończeniu jego importowania przycisk *Next* (Dalej). Jednocześnie zostanie wyświetlony monit przypominający o wybraniu typu pliku tekstowego, który pojawi się po lewej stronie w formie rozwijanej listy. Program sam próbuje określić typ wybranego pliku tak, aby m.in. prawidłowo wyświetlać znaki w tekście. Przy czym, jeśli domyślnie zostanie wczytany nieprawidłowy typ pliku, wówczas należy dokonać korekty polegającej na zmianie ustawień tak, aby pasował on do specyfiki dokumentu źródłowego. Aby mieć pewność, że operacja dotycząca wyboru właściwego typu zapisu pliku przebiegła prawidłowo, program wyświetla w trybie podglądu fragment jego treści. Kiedy upewnimy się, że wszystko jest w porządku, należy kliknąć na przycisk *Next* (Dalej).

Następnie użytkownik ma możliwość wyboru języka tożsamego z językiem zaimportowanego dokumentu tekstowego. Niezależnie od tego program próbuje automatycznie dopasować język, a także dostosować opcje związane z listą słów (*Word list*). W ten sposób można na przykład zdecydować, że niektóre ciągi znaków będą traktowane jako pojedyncze lub odrębne słowa. Te wybory mogą być ważne w przypadku korzystania z kilku automatycznych funkcji CATMA, takich jak funkcje przeliczania lub przeszukiwania (o czym będzie jeszcze mowa). Z technicznego punktu widzenia, aby dodać ciąg kreślonych znaków, należy go wpisać w polu u dołu okna i kliknąć przycisk *Add entry* (Dodaj wpis), a następnie przycisk *Save list* (Zapisz listę) i kliknąć *Next* (Dalej).



Ilustracja 2.2. Okno Menedżera programu CATMA

Źródło: opracowanie własne na podstawie programu CATMA

Add new Source Document

1. Source Document location

2. Source Document file type

3. Wordlist options

4. Content details

Documents

File Name	Title	Author	Description	Publisher
W_2.doc	W 2			

Cancel

Back

Next

Finish

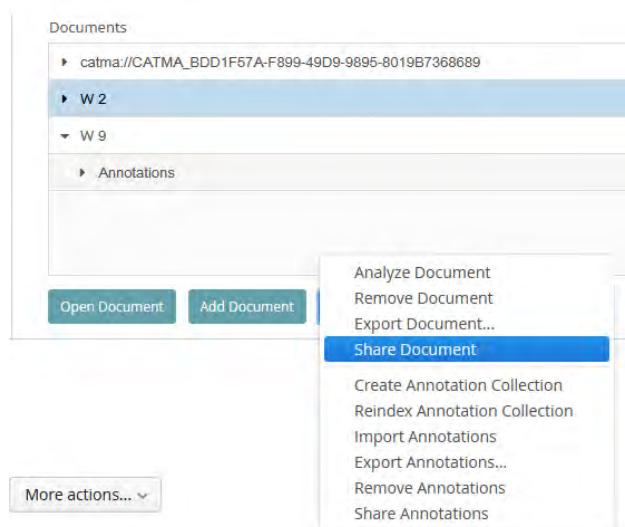
Ilustracja 2.3. Dodawanie dokumentu tekstowego do sekcji *Documents*

Źródło: opracowanie własne na podstawie programu CATMA

Teraz pozostaje jeszcze wprowadzić szczegółowe informacje o tekście, np. jego nazwę i autora. Natomiast, aby ostatecznie dodać tekst, należy kliknąć na przycisk *Finish* (Zakończ). Od tego momentu dokument tekstowy znajduje się w sekcji *Documents* (Dokumenty) w *Repository Manager* (Menedżera repozytoriów).

Uwaga: Podczas dodawania pierwszego dokumentu źródłowego CATMA generuje zestaw przykładowych elementów, takich jak kolekcja znaczników użytkownika, aby przechowywać znaczniki i Bibliotekę znaczników z przykładowym zestawem tagów.

Program daje też możliwość udostępniania dokumentu źródłowego innym osobom. Aby to uczynić, należy wybierać dokument, który chce się udostępnić, kliknąć na przycisk *More actions...* (Więcej działań...), a następnie *Share document* (Podziel dokument). W dalszej kolejności trzeba wprowadzić adres e-mail użytkownika, któremu chce się udostępnić dokument, oraz określić, czy ma on uprawnienia do edytowania dokumentu (opcja: *write*), czy jedynie do jego odczytu (opcja: *read*). Następnie kliknij przycisk *Save* (Zapisz).



Ilustracja 2.4. Funkcja udostępnienia dokumentu tekstowego innym osobom

Źródło: opracowanie własne na podstawie programu CATMA

Jeśli dany dokument jest współdzielony z inną osobą, wówczas, aby miała ona możliwość śledzenia bieżących zmian, dokonywanych w jego tekście, należy co jakiś czas kliknąć na przycisk *Refresh* (Odśwież) znajdujący się w prawym górnym rogu

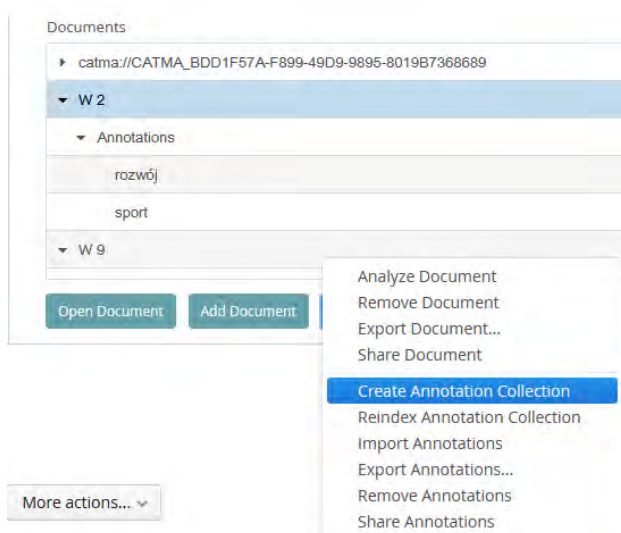
okna programu. Opcja ta działa w dwie strony, to znaczy modyfikacje dokonywane przez osobę, której udzielono uprawnień do edycji dokumentu, będą widoczne, gdy naciśnie ona ten sam przycisk w oknie programu zainstalowanym na jej komputerze.

2.1.2. Eksportowanie tekstów

Możliwe jest również eksportowanie dokumentów z CATMA do komputera, np. w celu utworzenia kopii zapasowej. Aby to zrobić, należy wybrać dany dokument z listy Repozytorium Menedżera i kliknąć na przycisk *More actions...* (Więcej działań...), a następnie wybierać opcję *Export document...* (Eksportuj dokument) i kliknąć przycisk *Save* (Zapisz).

2.2. Tworzenie adnotacji

Jeśli chcemy utworzyć własną kolekcję adnotacji, należy zaznaczyć wybrany tekst, znajdujący w sekcji *Documents*. Następnie klikamy na przycisk *More actions...* (Więcej działań...) w sekcji *Documents* i wybieramy opcję *Create Annotation Collection* (Utwórz kolekcję adnotacji). Teraz należy wpisać nazwę adnotacji i kliknąć przycisk *Save* (Zapisz). Wszystkie utworzone w ten sposób znaczniki zostaną wyświetlone po kliknięciu małej strzałki przed wybranym dokumentem w sekcji *Documents*.



Ilustracja 2.5. Tworzenie adnotacji przez użytkownika

Źródło: opracowanie własne na podstawie programu CATMA

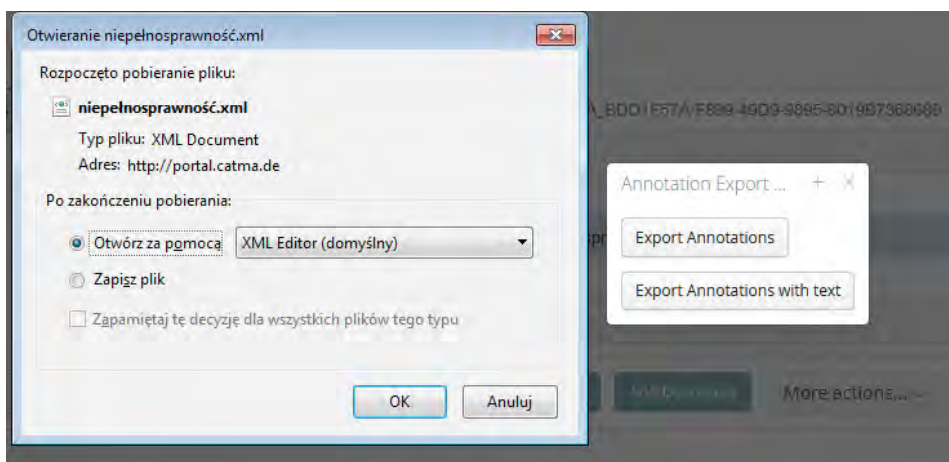
Uwaga: Jeśli dodamy tekst zaznaczony zgodnie ze standardami XML, kolekcja adnotacji zostanie automatycznie zaimportowana do CATMA jako kolekcja statyczna.

2.2.1. Dzielenie się adnotacjami

Jeśli chcemy podzielić się z kimś kolekcją adnotacji, należy wybrać daną kolekcję, a następnie kliknąć na przycisk *More actions...* i *Share Annotations*. Teraz trzeba wpisać adres e-mail użytkownika, z którym chcemy podzielić się kolekcją adnotacji, i ustalić, czy będzie on miał uprawnienia do jej edytowania, czy wyłącznie do odczytu. Funkcja ta jest analogiczna do udostępniania dokumentów, podobnie jak aktualizacja dokonywanych zmian przez któregokolwiek z użytkowników poprzez klikanie na przycisk *Refresh* (Odśwież).

2.2.2. Eksportowanie i importowanie kolekcji adnotacji

Program pozwala również na eksportowanie kolekcji adnotacji z CATMA do komputera w celu utworzenia kopii zapasowych. Czynność ta wykonywana jest w sposób tożsamy dla eksportowania dokumentów.



Ilustracja 2.6. Eksportowanie kolekcji adnotacji

Źródło: opracowanie własne na podstawie programu CATMA

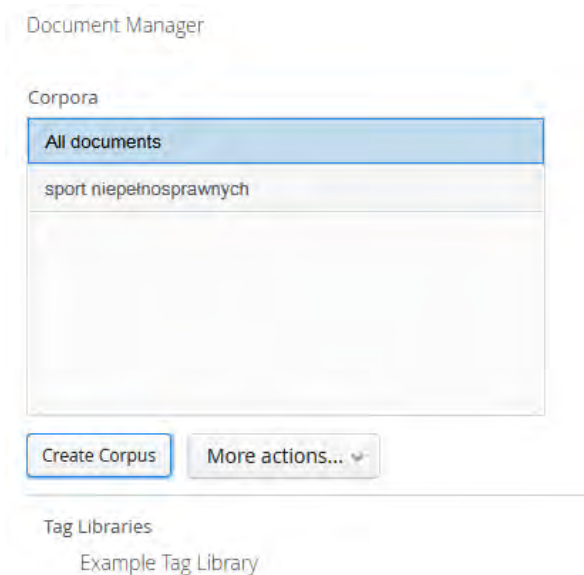
Aby ponownie zaimportować kolekcję adnotacji, należy zaznaczyć Dokument, do którego chcemy, aby była załączona, a następnie kliknąć na przycisk *More*

actions... i wybrać opcję *Import Annotations* (Importuj adnotacje). Teraz trzeba jeszcze kliknąć na przycisk Wybierz plik i dwukrotnie kliknąć na kolekcję adnotacji, którą chce się zaimportować.

Uwaga: CATMA jest wstecznie kompatybilny w tym sensie, że wszystkie adnotacje utworzone przy użyciu wcześniejszych wersji są importowalne.

2.3. Tworzenie korpusu dokumentów

Program pozwala na zestawianie importowanych dokumentów w korpusy (*Corpora*). Domyślnie wszystkie dodawane dokumenty zostają włączone do korpusu *All documents*. Użytkownik ma jednak możliwość tworzenia własnych zbiorów danych (korpusów dokumentów). Aby to uczynić, wystarczy kliknąć na przycisk *Create Corpus* znajdujący w oknie programu w sekcji *Corpora*, a następnie wprowadzić jego nazwę i kliknąć *Save* (Zapisz).



Ilustracja 2.7. Tworzenie korpusu dokumentów

Źródło: opracowanie własne na podstawie programu CATMA

Dodawanie poszczególnych dokumentów do dowolnego korpusu polega na ich przeciągnięciu i upuszczeniu we właściwym miejscu. Przy czym warto pamiętać

tać, że każdy dokument dołączony do określonego korpusu jest również częścią korpusu *All documents*, bowiem nowo powstałe korpusy są automatycznie do niego „podpinane” (stają się swego rodzaju podzbiorami zbioru obejmującego wszystkie dokumenty).

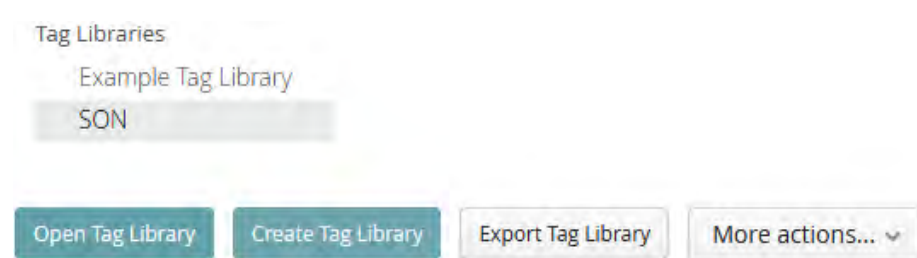
I podobnie jak miało to już miejsce w przypadku kolekcji znaczników czy poszczególnych dokumentów, także korpusy można udostępniać innym osobom, określając parametry ich uprawnień do edytowania lub wyłącznie ich odczytywania. Również bieżąca aktualizacja korpusu przez dowolnego użytkownika odbywa się na zasadzie odświeżania jego zawartości.

2.4. Znaczniki (tagi)

Znaczniki, a więc tagi to kategorie, które stosuje się do tekstu w celu jego analizy. Tagi przechowywane są w strukturze hierarchicznej w ramach *Tag Manager* (Menedżera tagów). Części tekstu opatrzone określonym znacznikiem są wyróżnione kolorem tagu. Natomiast wszystkie czynności związane z „tagowaniem” są wykonywane przy użyciu odpowiedniego modułu.

2.4.1. Tworzenie Biblioteki znaczników

Aby utworzyć własny zestaw znaczników, należy kliknąć na przycisk *Create Tag Library* w sekcji Tag Library, która znajduje się w lewym dolnym rogu okna programu. Następnie wystarczy wpisać nazwę nowo powstającej Biblioteki znaczników i nacisnąć przycisk *Save* (Zapisz).



Ilustracja 2.8. Tworzenie Biblioteki znaczników (tagów)

Źródło: opracowanie własne na podstawie programu CATMA

Biblioteki znaczników można udostępniać innym osobom, określając parametry ich uprawnień do edytowania lub wyłącznie ich odczytywania. Również bieżąca

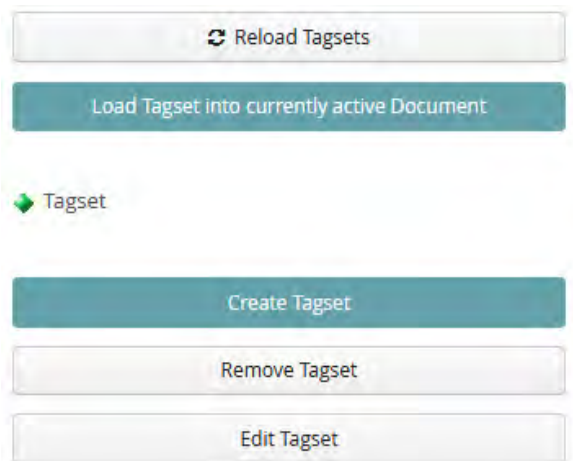
aktualizacja biblioteki przez dowolnego użytkownika odbywa się na zasadzie odświeżania jego zawartości, a więc tak jak zostało to już wcześniej wskazane przy okazji opisywania dokumentów i ich korpusów.

2.4.2. Tworzenie znaczników

Aby utworzyć własne zestawy tagów, należy wybierać Bibliotekę znaczników z sekcji Tag Library i kliknąć na przycisk *Open Tag Library* w celu otwarcia Menedżera tagów.

Uwaga: Jeśli otworzysz więcej niż jedną Tag Library w jednej sesji roboczej, poszczególne biblioteki będą wyświetlane w różnych kartach w oknie Menedżera tagów. Jeśli okno zostanie zamknięte i ponownie otwarte w jednej sesji roboczej, wszystkie zakładki, które były poprzednio otwarte, zostaną ponownie wyświetlone.

W lewej części okna Tag Manager, czy inaczej mówiąc „Taggera”, wyświetlany jest zestaw znaczników, poszczególne tagi, a także przypisane im kolory oraz ich właściwości. Część prawa okna używana jest do tworzenia, usuwania i edytowania zestawów tagów, znaczników i ich właściwości.



Ilustracja 2.9. Okno Tag Manager

Źródło: opracowanie własne na podstawie programu CATMA

Należy pamiętać, że poszczególne tagi można tworzyć tylko wówczas, gdy utworzy się Bibliotekę tagów (*Tagset*). Aby to uczynić, należy kliknąć na przycisk *Create Tagset*

SON x

Tags

Tag Color

▼ sport osób niepełnosprawnych	
znięczenie	
wysiek	
sukces	
porazka	
zmiana	
rodzina	
wsparcie	
trudności	
zaangażowanie	
regulaność	
wyrwnałość	
akceptacja	
przyjaźń	
wywalizacja	

Reload Tagsets

Load Tagset into currently active Document

Tagset

Create Tagset

Remove Tagset

Edit Tagset

Tag

Create Tag

Remove Tag

Edit Tag

Ilustracja 2.10. Tworzenie nowego tagu

Źródło: opracowanie własne na podstawie programu CATMA

i wpisać nazwę nowo tworzonej Biblioteki tagów. Utworzony zestaw tagów jest wyświetlany w sekcji *Tagset* Menedżera tagów. Teraz można już tworzyć poszczególne tagi, co odbywa się poprzez kliknięcie na przycisk *Create Tag*. Tworząc nowy tag, należy wprowadzić jego nazwę, wybrać przypisany mu kolor oraz potwierdzić wskazane czynności, klikając na przycisk *Save* (Zapisz). Wybrany kolor zostanie użyty do zaznaczenia sekcji oznaczonych w tekście. Utworzony tag będzie zaś wyświetlony w sekcji tagi w Menedżerze tagów, po wcześniejszym kliknięciu na małą strzałkę znajdującą się przed właściwym tagiem.

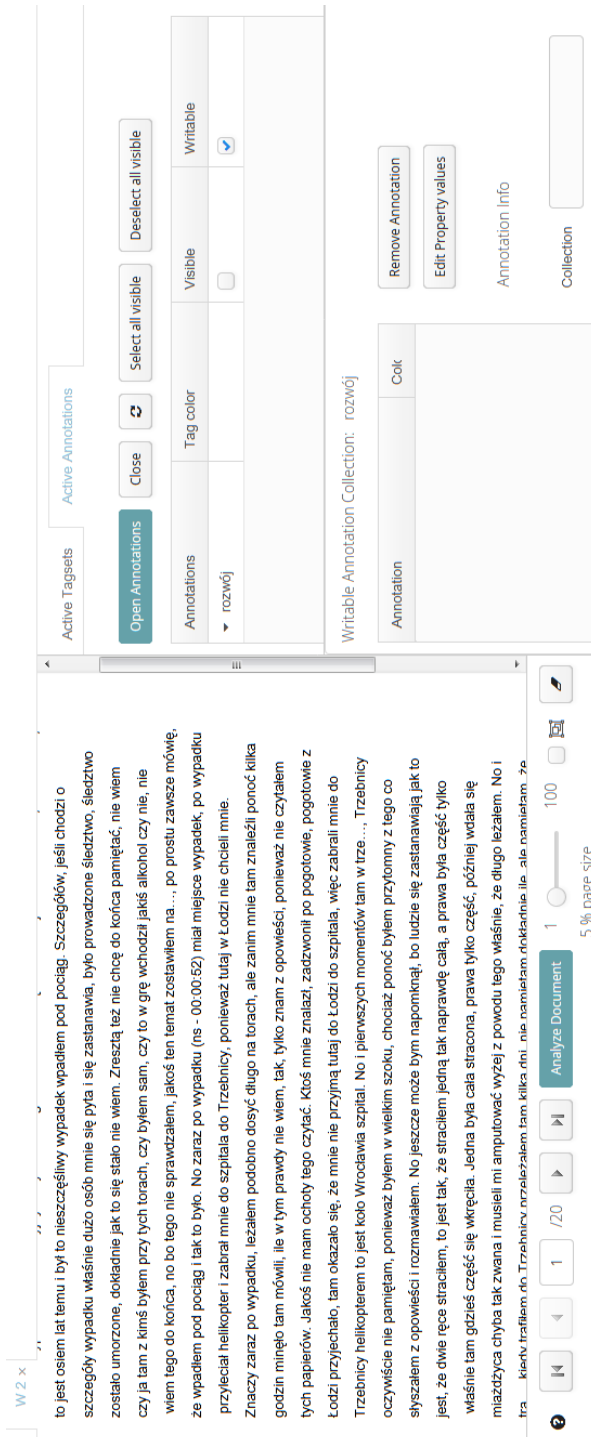
W programie CATMA można również utworzyć znaczniki podrzędne dla istniejących już tagów znajdujących się w Menedżerze tagów. Do tego celu służy opcja *Create Tag*. Znaczniki utworzone w ten sposób zostaną automatycznie przypisane jako podrzędne do wcześniej wybranego tagu.

2.4.3 Oznaczanie tekstu tagami (tagowanie)

Aby rozpocząć rzeczywiste ‘tagowanie’, a więc proces polegający na przypisywaniu znaczników do wybranych fragmentów dokumentów tekstowych, należy kliknąć na małą strzałkę znajdującą się przed tekstem, który chcemy ‘otagować’ w sekcji *Documents*. Następnie należy kliknąć na strzałkę przed kolekcjami znaczników (*Annotations*) i z rozwiniętej listy wybierać konkretną kolekcję znaczników. Teraz wystarczy już tylko kliknąć na opcję *Open Annotations*, co spowoduje otwarcie okna ‘Taggera’.

Okno Taggera (ilustracja 2.11) składa się z kilku obszarów. Największy z obszarów zajmuje ten, w którym wyświetlana jest treść wybranego dokumentu tekstowego (po lewej stronie okna). Warto zaznaczyć, że ułatwieniem dla użytkownika pracującego w programie jest możliwość regulacji wielkości czcionki, a więc ilość wyświetlanego tekstu na stronie. Po prawej stronie okna Taggera znajdują się aktywne tagi oraz zbioru adnotacji. W zależności od tego, jaką zakładkę wybierzemy, wyświetlane są kolekcje adnotacji i adnotacje, które zostały użyte do oznaczenia wybranych fragmentów tekstu lub tagi. Jeśli klikniemy na część tekstu, która została otagowana, wówczas spowoduje to wyświetlenie się użytych do tego celu tagów. Dzięki temu możemy również dodawać kolejne tagi bądź usuwać te, które uznamy za niepotrzebne.

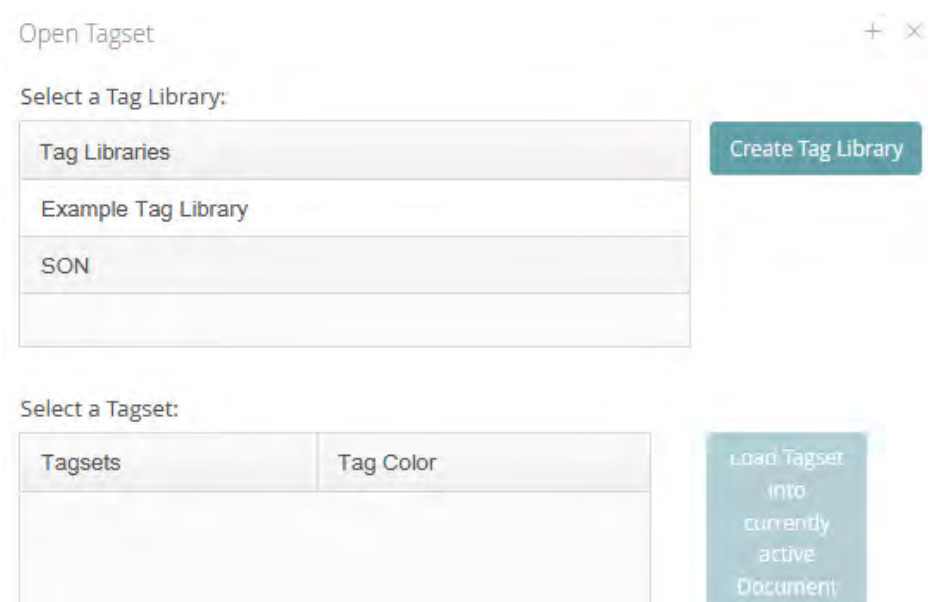
Uwaga: jeśli w jednej sesji roboczej zostanie otwartych więcej niż jeden dokument, będą one wyświetlane w różnych kartach w oknie Taggera. Jeżeli okno zostanie zamknięte i ponownie otwarte w tej samej sesji roboczej (za pośrednictwem funkcji dostępnej na pasku menu lub za pomocą przycisku *Open Document / Open Markup Collection* (Otwórz dokument / Otwórz znacznik kolekcji)), wówczas wszystkie karty, które były już otwarte, wyświetlą się ponownie.



Ilustracja 2.11. Wygląd okna 'Taggera'

Źródło: opracowanie własne na podstawie programu CATMA

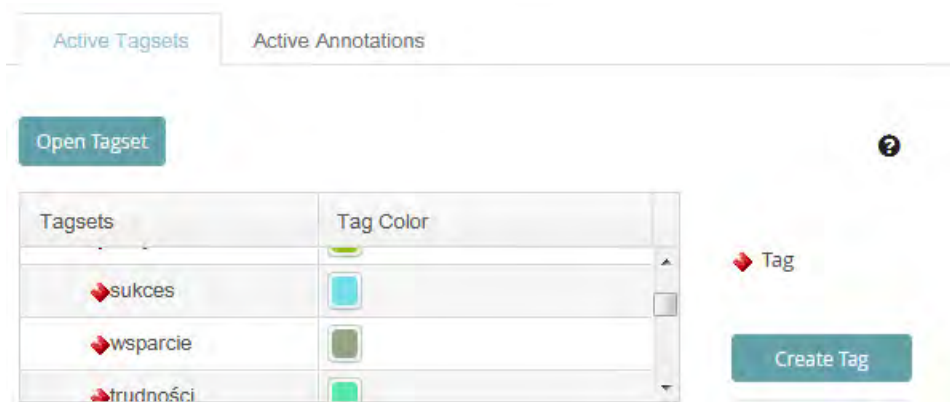
Przy czym, aby faktycznie móc opatrywać kolejnymi znacznikami (tagami) wyświetlane dokumenty tekstowe, musimy oprócz okna Taggera otworzyć także Bibliotekę znaczników (*Tag Library*) zawierającą znaczniki (tagi), których chcemy używać do oznaczania tekstu. W tym celu należy wybrać konkretną Bibliotekę znaczników z sekcji *Tag Library* w Menedżerze repozytoriów (*Repository Manager*) i kliknąć na przycisk *Open Tag Library*, aby otworzyć okno Menedżera tagów (*Tag Manager*). Teraz wystarczy już tylko z Biblioteki tagów (*Tag Library*) przeciągnąć i upuścić zestaw tagów (*Tagset*), z którego chcemy korzystać, do sekcji *Active Tagsets* znajdującej się w Taggerze.



Ilustracja 2.12. Wybieranie zestawu tagów, które chcemy wykorzystywać w procesie tagowania tekstu

Źródło: opracowanie własne na podstawie programu CATMA

Następną czynnością jest kliknięcie na strzałkę znajdującą się przed *Tagset* w oknie Taggera w celu wyświetlenia utworzonych przez nas zestawów tagów. Teraz klikamy na kartę aktywnych znaczników Taggera i zaznaczamy kolekcję znaczników, w której chcemy zapisać swoje prace (jeśli nie jest jeszcze wyświetlona, klikamy na strzałkę przed kolekcjonerem Markup użytkownika w kolumnie, gdzie dane te będą zapisywane). Wreszcie należy wybrać fragmenty tekstu znajdujące się po lewej stronie okna Taggera, a następnie kliknąć przycisk *Tag Color* na karcie *Active Tagsets* po prawej stronie Taggera, aby zaznaczyć i opatrzyć tagami odpowiednie fragmenty tekstu.



Ilustracja 2.13. Oznaczanie i tagowanie wybranych fragmentów tekstu

Źródło: opracowanie własne na podstawie programu CATMA

Jeśli chcielibyśmy oznaczyć danym tagiem tekst podzielony na dwa fragmenty, wówczas należy wybrać pierwszą część takiego tekstu, a następnie przytrzymać klawisz Ctrl na klawiaturze i wybierać kolejną część tekstu. Następnie kliknąć przycisk Tag na karcie Active Tagsets znajdujący się po prawej stronie Taggera.

Jeśli chcemy zamknąć zestaw tagów w sekcji Active Tagsets lub w kolekcji użytkownika w sekcji Active Annotations, należy kliknąć prawym przyciskiem myszy ikonę *Tagset* lub *Annotations Collection* i wybrać przycisk *Close*.

2.4.4. Edycja i usuwanie znaczników

Aby edytować dany tag, wybierz go w Menedżerze tagów lub jeśli już został przeciągnięty i upuszczony do Taggera, w jego oknie kliknij na przycisk *Edit Tag* (Edytuj tag) i zmień jego nazwę lub kolor. Aby usunąć tag wraz ze wszystkimi wystąpieniami tagu w tekście, wybierz go w Menedżerze tagów lub Taggerze i kliknij na przycisk *Remove Tag* (Usuń tag). Jeśli chcesz usunąć z tekstu tylko oznaczenie tagiem danego fragmentu tekstu, kliknij na tę część tekstu, która wyświetlona jest po lewej stronie okna Taggera, a następnie wybierz odpowiedni tag, który chcesz usunąć z listy znajdującej się w prawej części Taggera, i kliknij opcję *Remove Tag Instance*.

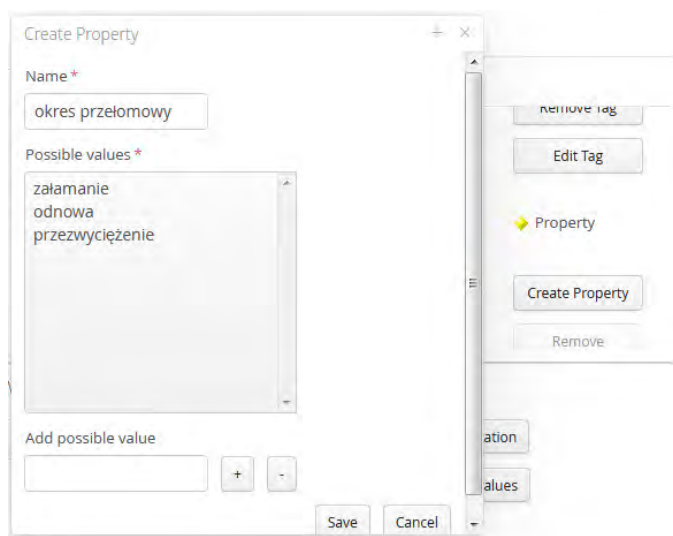
Aby edytować nazwę całego zestawu tagów (Tagset), zaznacz tagi w Menedżerze tagów lub Taggerze i kliknij *Edit tag*. Możesz również edytować zestaw tagów, dodając lub usuwając tagi w Menedżerze tagów lub Taggerze. Każda zmiana jest automatycznie wdrażana w obu miejscach. Aby usunąć tagi, zaznacz zestaw tagów i kliknij opcję *Remove Tag* (Usuń tag).

2.5. Właściwości

Program pozwala również na przypisywanie określonych właściwości do znaczników (tagów), co w efekcie umożliwia łatwiejsze ich odróżnienie oraz udoskonalenie m.in. dotyczące konkretyzacji znaczenia. Właściwości pozwalają również na tworzenie rozbudowanych poziomych struktur składających się z zestawów tagów (Tagsets), a także pionowej struktury budowanej za pośrednictwem podrzędnych znaczników (Subtags). Dzieje się tak, bowiem można wybrać wiele wartości dla właściwości przypisywanych do danego tagu.

2.5.1. Tworzenie właściwości

Aby dodać jakąś właściwość, należy otworzyć okno Taggera lub Menedżera tagów za pośrednictwem paska menu i wybrać żądany tag. Jeśli nie są wyświetlane żadne tagi, trzeba najpierw kliknąć na strzałkę znajdującą się przed Tagset. Teraz możemy już kliknąć na przycisk *Create Property* i wpisać nazwę właściwości w wyskakującym oknie. Aby dodać wartość, trzeba ją wpisać w polu u dołu okna podręcznego i kliknąć przycisk plusa (+). Wartość jest teraz wyświetlana na liście. Jeśli chcesz dodać więcej wartości, postępujemy w ten sam sposób. Jeśli chcemy usunąć wartość z listy, trzeba ją wybrać i kliknąć odpowiedni przycisk. Po zakończeniu klikamy na przycisk *Save*.



Ilustracja 2.15. Edytowanie właściwości przypisywanych do tagu

Źródło: opracowanie własne na podstawie programu CATMA

2.5.2. Praca z właściwościami

Jeśli chcesz umieścić w tekście wartość właściwości na potrzeby danego tagu, należy otworzyć kolekcję znaczników, a następnie przeciągnąć i upuścić zestawy tagów (Tagsets) zawierające tag z żądanymi wartościami właściwości z Menedżera tagów (Tag Manager) do sekcji Active w oknie Taggera. Jeśli w zestawie tagów (Tagset) nie ma jeszcze określonych właściwości, można je bezpośrednio utworzyć w Active Tagsets. Aby to uczynić, należy kliknąć na wybrany fragment tekstu, który zawiera tag, do którego to chcemy przypisać jakąś wartość właściwości. Tag jest teraz wyświetlany w prawym dolnym rogu Taggera. Teraz kliknij strzałkę przed tagiem i wybierz jego właściwość. Kliknij przycisk *Edit properties*, wybierz jedną lub więcej wartości z listy i kliknij przycisk *Save*.

2.6. Wznawianie pracy nad tekstem

Oczywiście program pozwala na elastyczną formę pracy użytkownika, co oznacza, że poza istniejącymi opcjami (o których już była mowa), można w każdej chwili zakończyć i powrócić do jej kontynuowania bez obawy o utratę danych. Jeśli zatem chcemy wznowić pracę nad dokumentem, który został już otwarty w poprzedniej sesji roboczej, wówczas należy wybrać odpowiednią kolekcję znaczników (*Annotations Collection*) z sekcji *Documents* znajdującą się w *Repository Manager* (Menedżerze repozytoriów). W sytuacji gdy nie jest wyświetlana żadna kolekcja znaczników, trzeba kliknąć na małą strzałkę znajdującą się przed odpowiednim dokumentem, a następnie po raz drugi, gdy taka strzałka pojawi się przed zbiorem znaczników użytkownika (*Annotations Collection*). Teraz wystarczy kliknąć na przycisk otwierania kolekcji znaczników (*Open Annotations*), aby otworzyć okno Taggera, w którym wyświetlone zostaną teksty i aktywne kolekcje znaczników. Jeśli zaś chcemy wyświetlić poprzednio zamieszczone tagi w tekście, trzeba zaznaczyć pole widoczne za zestawami tagów (*Tagsets*) w sekcji *Active Markup Collections* okna Taggera. Przy czym użytkownik może także wyświetlić jeden lub kilka wybranych tagów w tekście, co jest szczególnie pomoce, gdy mamy już sporą ilość znaczników, a chcemy skoncentrować się tylko na pewnej ich części.

Jeśli natomiast chcemy wznowić proces tagowania, wówczas musimy wybrać kartę *Active Tagsets* w Taggerze.

Uwaga: Znaczniki używane do oznaczania tekstu w poprzedniej sesji roboczej nie są automatycznie wyświetlane w Taggerze. Dzieje się tak, ponieważ modyfikacji uległ zestaw tagów używanych w Menedżerze tagów, co nastąpiło już po otagowaniu tekstu. Ma to chronić pracę użytkownika, bowiem jeśli

modyfikowane tagi były automatycznie importowane do Taggera, wówczas również kolekcja znaczników użytkownika (*User Markup Collection*) zostałaby automatycznie zmieniona, co uniemożliwiłoby przejrzanie poprzednich wyników pracy.

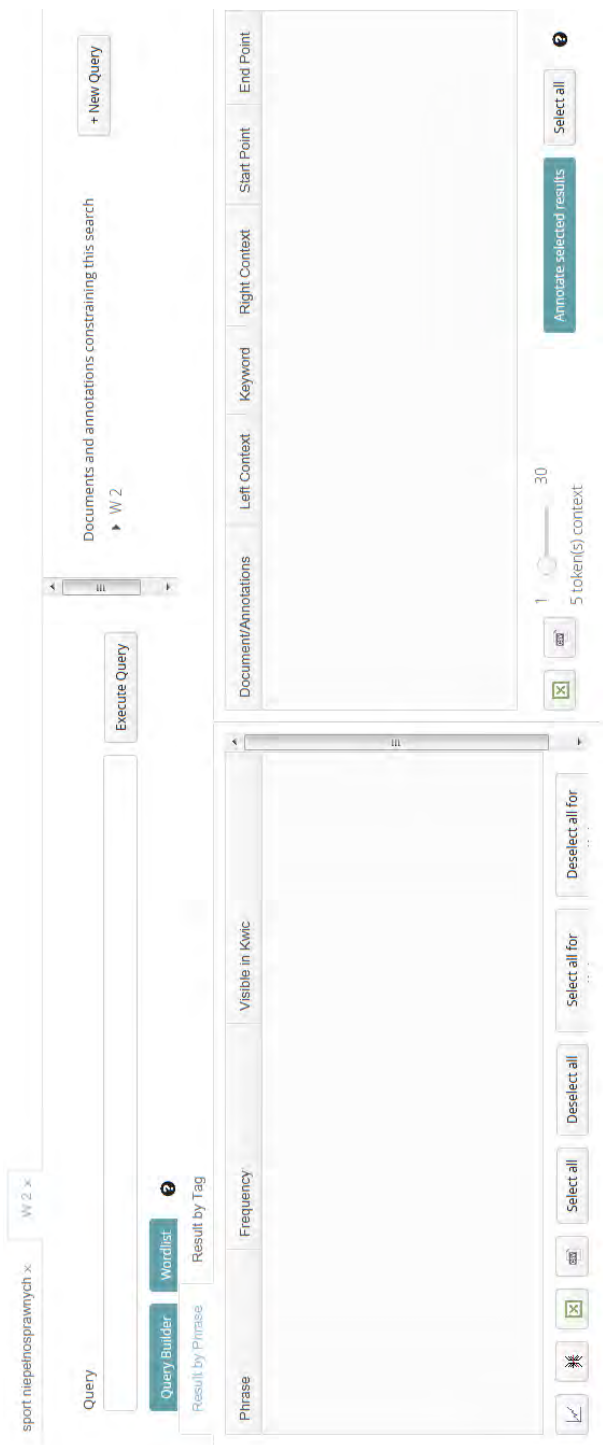
Aby zaimportować zestaw tagów (*Tagset*), należy otworzyć menedżer tagów, wybierając odpowiednią Bibliotekę tagów (*Tag Library*) spośród znajdujących się po lewej stronie Menedżera repozytoriów i kliknąć na przycisk *Open Tag Library*. Teraz należy przeciągnąć i upuścić żądany zestaw tagów z Menedżera tagów (*Tag Manager*) do sekcji *Active Tagsets* znajdującej się w oknie Taggera. Jeśli zmodyfikowano zestaw tagów po ostatniej operacji tagowania, zostanie wyświetlone pytanie, czy chcemy zaktualizować załączoną kolekcję znaczników. Jeśli tak, to wówczas klikamy na przycisk OK.

2.7. Analiza tekstu

Moduł *Analizer* pozwala na zastosowanie różnych funkcji analizy tekstu. Może to być na przykład wygenerowanie listy słów w tekście zestawionych w porządku alfabetycznym lub według częstotliwości ich występowania. W połączeniu z innymi modułami można również użyć funkcji analizy, m.in. do wyszukiwania słów lub fraz w tekstach, a następnie oznaczania (tagowania) wyników wyszukiwania, a także generowania listy wszystkich miejsc, w których użyto danego tagu w tekście, co z kolei stanowi podstawę do tworzenia mapy rozkładu tagów tworzonej za pomocą *Visualizera*.

Aby użyć dowolnej funkcji analizy, należy otworzyć okno *Manage Resources* i kliknąć przycisk *Analyze Document* pod wyświetlonym tekstem (aby pokazała się jego treść, wystarczy kliknąć na jego nazwę w obszarze *Documents*), co spowoduje pojawienie się okna Analizatora.

Okno Analizera (ilustracja 2.17) podzielone jest na kilka sekcji. W lewym górnym rogu można uruchomić listę słów kluczowych tzw. *Wordlist* lub wykonać kwerendę za pomocą narzędzia *Query Builder*, bezpośrednio wpisując zapytanie. Z kolei w prawym górnym rogu znajduje się przegląd różnych dokumentów źródłowych i kolekcji znaczników, które są uwzględnione w bieżącej analizie. W lewej dolnej części wyświetlana jest lista słów lub kwerend wraz z ich częstotliwościami występowania. Natomiast w prawej dolnej części okna widoczne są listy wybranych słów kluczowych w kontekstach ich użycia w tekście.

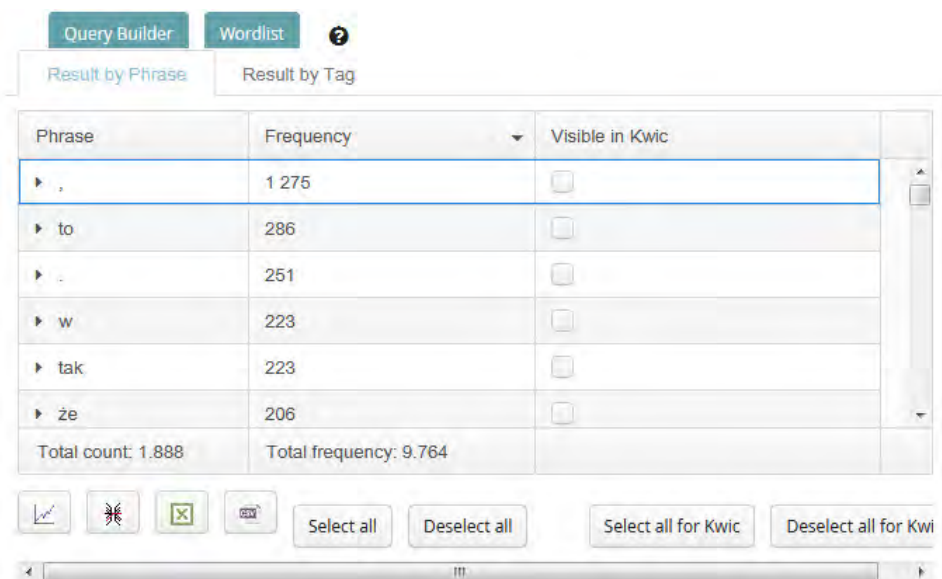


Ilustracja 2.17. Wygląd okna Analizera w programie CATMA

Źródło: opracowanie własne na podstawie programu CATMA

2.7.1. Tworzenie list słownictwa

Aby wyświetlić listę słów kluczowych, należy kliknąć przycisk *Wordlist* w górnym lewym rogu *Analizer*. Lista słów pojawi się w lewym dolnym rogu okna Analizatora. Obok każdego ze słów widoczna jest także częstotliwość jego występowania w tekście. Pod listą po lewej stronie można natomiast znaleźć liczbę typów słów (liczba całkowita), a po prawej stronie – częstotliwość występowania poszczególnych wyrazów (całkowita częstotliwość).



The screenshot shows the 'Wordlist' tab in the CATMA software. At the top, there are buttons for 'Query Builder', 'Wordlist', and a help icon. Below these are two tabs: 'Result by Phrase' (selected) and 'Result by Tag'. The main area contains a table with three columns: 'Phrase', 'Frequency', and 'Visible in Kwic'. The table lists several phrases with their corresponding frequencies. At the bottom of the table, it shows 'Total count: 1.888' and 'Total frequency: 9.764'. Below the table, there are several icons (a line graph, a spider, a green X, and a document) and four buttons: 'Select all', 'Deselect all', 'Select all for Kwic', and 'Deselect all for Kwic'.

Phrase	Frequency	Visible in Kwic
,	1 275	☐
to	286	☐
.	251	☐
w	223	☐
tak	223	☐
że	206	☐
Total count: 1.888		Total frequency: 9.764

Ilustracja 2.18. Widok listy słów kluczowych w oknie Analizera

Źródło: opracowanie własne na podstawie programu CATMA

Można również wybrać kryterium, według którego lista słów ma być konstruowana. Przykładowo, jeśli chcemy, aby wyrazy były wyświetlane w kolejności alfabetycznej, powinniśmy kliknąć napis *Phrase* znajdujący się na górze pierwszej kolumny listy słów. Przy napisie *Phrase* pojawi się mała strzałka skierowana w górę, a lista słówek zostanie wygenerowana w kolejności alfabetycznej. Jeśli chcemy odwrócić kolejność, powinniśmy ponownie kliknąć na *Phrase*. Grot strzałki zostanie skierowany w dół, a lista słów wyświetlana jest w odwrotnym porządku alfabetycznym. Możemy także zdecydować się na sporządzenie listy zgodnie z częstotliwością występowania fraz w tekście. Aby to zrobić, klikamy na napis *Frequency* znajdujący się na górze drugiej kolumny listy słów. Spowoduje to wyświetlenie listy według częstotliwości od najczęściej do najrzadziej występujących słów.

2.7.2. Bieżące zapytania

Moduł Analizera oferuje również możliwość bezpośredniego wyszukiwania słów, części słów, kombinacji wyrazów lub znaczników w tekście. Wyniki tych zapytań mogą być następnie wykorzystywane dla dalszych operacji, takich jak tagowanie lub wizualizacja (o czym będzie jeszcze mowa później). Warto zaznaczyć, że istnieją dwa różne sposoby uruchamiania zapytań: można bezpośrednio wpisać zapytania lub użyć specjalnego kreatora zapytania. Należy pamiętać, że dla zapytań wymagane jest wykorzystanie specjalnych operatorów i określonej struktury składniowej, co zostało opisane poniżej. Łatwiejsze może się jednak okazać użycie konstruktora kwerend tzw. *Query Builder*, opartego na języku naturalnym, który pomaga krok po kroku i intuicyjnie tworzyć zapytania i przekształcać kwerendę tak, aby spełniała wymagania badacza. Dlatego temu właśnie narzędziu przyjrzymy się w pierwszej kolejności.

2.7.3. Narzędzie przeprowadzania kwerend

Program CATMA został wyposażony w rozbudowany system przeszukiwania danych. Narzędzia, jakie ma w tej kwestii do dyspozycji użytkownik, zostały podzielone na dwie główne kategorie. Do pierwszej zalicza się zapytania wykonywane za pomocą kreatora kwerend. Druga kategoria reprezentuje zaś bardziej zaawansowane, ale też trudniejsze do zastosowania opcje przeszukiwania, wymagające od użytkownika pewnej wprawy, dotyczącej podstaw programowania. Poniżej w sposób syntetyczny staram się opisać oba sposoby przeprowadzania kwerend.

2.7.3.1. Proste zapytania z zastosowaniem kreatora kwerend

Aby użyć Kreatora zapytań, należy otworzyć *Analizyzer*, klikając przycisk *Analyze Document* z listy *More actions...* znajdującej się w oknie *Manage Resources* w sekcji *Documents*. Po otwarciu okna Analizatora klikamy przycisk *Query Builder* w lewym górnym rogu okna, co powoduje wyświetlenie się kreatora.

Warto zaznaczyć, że program daje nam wybór, możemy bowiem zdecydować, jak chcemy prowadzić wyszukiwanie: według słów lub fraz, według stopnia podobieństwa, według tagów, kolokacji lub częstotliwości.

2.7.3.2. Wyszukiwanie przez słowo lub wyrażenie

Jedną z opcji jest wyszukiwanie słów lub fraz (*by word or phrase*). W tym celu należy wybrać słowo lub frazę i kliknąć przycisk Dalej. Możemy również zdecydować, czy chcemy wyszukać jedno dokładne słowo, czy też części słowa. Ta ostatnia

Query Builder

+

×

1. How do you want to search?

☒ by word or phrase

☐ by grade of similarity

☐ by Tag

☐ by collocation

☐ by frequency

Cancel

Back

Next

Finish

2. How does your phrase look like?

Ilustracja 2.19. Okno *Query Builder*

Źródło: opracowanie własne na podstawie programu CATMA

Query Builder

1. How do you want to search?

How does your phrase look like?

☒ by word or phrase
☐ by grade of similarity
☐ by tag
☐ by collocation
☐ by frequency

Cancel Back Next Finish

Ilustracja 2.20. Okno kreatora kwerend (Query Builder)

Źródło: opracowanie własne na podstawie programu CATMA

Query Builder

1. How do you want to search?

The first word starts with

niepeł

The first word contains

The first word ends with

The first word is exactly

2. How does your phrase look like?

Phrase	Frequency
▶ niepełnosprawnością	10
▶ niepełnosprawność	4
▶ niepełnosprawnych	2
▶ niepełnosprawnymi	1

☐ continue to build a complex query

Cancel

Back

Next

Finish

Ilustracja 2.21. Przykład przeszukiwania z wykorzystaniem wybranej frazy

Źródło: opracowanie własne na podstawie programu CATMA

opcja może być użyteczna, jeśli naszym zamiarem jest wyszukanie danego słowa, biorąc przy tym pod uwagę różne jego formy bądź odmiany. Do dyspozycji są zatem cztery różne możliwości: Pierwsze słowo zaczyna się od (*The first word starts with*), Pierwsze słowo zawiera (*The first word contains*), Pierwsze słowo kończy się na (*The first word ends with*), Pierwsze słowo jest dokładnie (*The first word is exactly*).

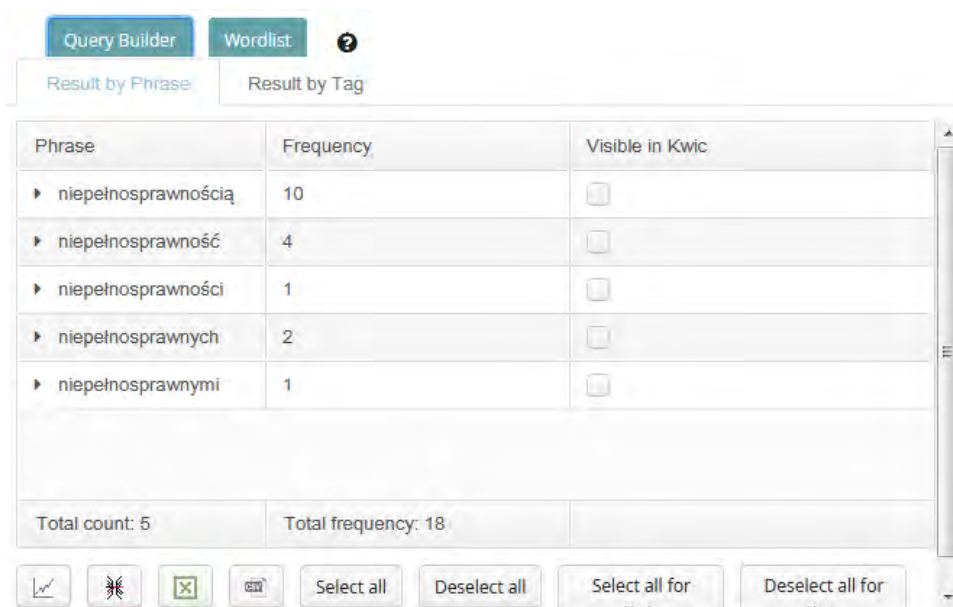
Aby rozpocząć proces przeszukiwania, należy wprowadzić daną frazę bądź słowo, które chce się wyszukać i kliknąć na przycisk *Show in preview*. Podglądy wyników wyszukiwania będą wyświetlane w dolnej części Kreatora.

Jeśli chcemy zmienić liczbę wyników wyświetlanych w podglądzie (*maximum total frequency*), wówczas należy wpisać żadaną maksymalną wartość w dolnej części Kreatora zapytań i kliknąć opcję *Show in preview*. Jeśli zdecydujemy się nie wprowadzać dokładnego słowa, ale skorzystamy z jednej z opcji przeszukiwania (np. Pierwsze słowo zaczyna się od) oraz gdy istnieją różne typy słów spełniające zapytania wśród wyników, wówczas słowa będą wyświetlane w różnych wierszach z jednoczesnym uwzględnieniem częstotliwości ich występowania. Na dole sekcji, po lewej stronie podglądu znajduje się całkowita liczba typów słów, które pasują do zapytania, zaś po prawej stronie całkowita częstotliwość wystąpień słowa w tekście.

Jeśli wyrażenie, którego szukamy, zawiera więcej niż jedno słowo, wówczas trzeba kliknąć przycisk dodania innego słowa (*add another word*). W ten sposób można wprowadzić kolejne słowo i ustalić jego pożądaną pozycję po pierwszym słowie. Na przykład: jeśli chcemy wyszukać dowolną formę czasownika „patrzeć”, wystarczy wpisać pierwsze słowo zaczynające się od „patrzeć”, a następnie kliknąć na *add another word* i dodać kolejne słowo „przed”, co oznacza, że znajduje się ono bezpośrednio po słowie wcześniejszym, a więc „patrzeć”. Jeśli zapytanie zawiera więcej niż dwa słowa, należy postępować w ten sam sposób z wpisywaniem następnych wyrazów.

Po zakończeniu klikamy przycisk *Finish*, co spowoduje wyświetlenie wyników kwerendy w oknie Analizatora. Wyniki są wyświetlane w taki sam sposób, jak ma to miejsce w podglądzie.

Uwaga: jeśli uruchomisz więcej niż jedną kwerendę w jednej sesji roboczej, wyniki wyszukiwania będą wyświetlane na różnych kartach okna Analizatora. Jeśli okno zostanie zamknięte i ponownie otwarte podczas tej samej sesji roboczej, wszystkie karty będą wyświetlane ponownie.



The screenshot shows the CATMA Query Builder interface. At the top, there are tabs for 'Query Builder' (selected) and 'Wordlist', along with a help icon. Below the tabs are two sub-tabs: 'Result by Phrase' (selected) and 'Result by Tag'. The main area displays a table with search results. The table has three columns: 'Phrase', 'Frequency', and 'Visible in Kwic'. There are five rows of results, each with a dropdown arrow in the 'Phrase' column. Below the table, there are summary statistics: 'Total count: 5' and 'Total frequency: 18'. At the bottom, there is a row of icons for various actions: a line graph, a crossed-out icon, a green 'X' icon, a 'copy' icon, and buttons for 'Select all', 'Deselect all', 'Select all for', and 'Deselect all for'.

Phrase	Frequency	Visible in Kwic
▶ niepełnosprawnością	10	☐
▶ niepełnosprawność	4	☐
▶ niepełnosprawności	1	☐
▶ niepełnosprawnych	2	☐
▶ niepełnosprawnymi	1	☐

Total count: 5 Total frequency: 18

Icons: [Line Graph] [Crossed Out] [Green X] [Copy] [Select all] [Deselect all] [Select all for] [Deselect all for]

Ilustracja 2.22. Wyświetlenie wyników przeszukiwania w oknie Analizatora

Źródło: opracowanie własne na podstawie programu CATMA

2.7.3.3. Przeszukiwanie według stopnia podobieństwa

Innym sposobem przeszukiwania, jaki można wykonać z zastosowaniem narzędzia zapytań, jest wyszukiwanie według stopnia podobieństwa. Aby skorzystać z tej funkcji, należy wybrać opcję *by grade of similarity* i kliknąć na przycisk *Next*. Następnie należy wprowadzić żądane słowo w lewym górnym rogu okna *Query Builder*. Wyniki wyszukiwania będą zawierać wprowadzone słowo, ale także słowa, które są w pewnym stopniu do niego podobne. Analizator używa w tym celu Algorytmu Rozpoznawania Wzorców. Można również określić pożądany stopień podobieństwa wyników wyszukiwania za pomocą skali *Grade of similarity*. Trzeba przy tym zaznaczyć, że skala działa według procentów. Na przykład: jeśli chcemy wyszukać słowa, które wykazują 70 procent podobieństwa do słowa „człowiek”, wówczas wpisujemy wskazane słowo i dostosowujemy stopień podobieństwa do 70. Warto przy okazji zwrócić uwagę, że jeśli wybierzemy 100 jako stopień podobieństwa, w wynikach zostanie wyświetlone tylko dokładne słowo, które wpisaliśmy jako podstawę przeszukiwania.

Aby wyświetlić podgląd bieżących wyników wyszukiwania, klikamy opcję *Show in preview*. Podglądy bieżących wyników wyszukiwania będą wyświetlane w dolnej części Kreatora zapytania.

Query Builder

1. How do you want to search?

2. The word is similar to

człowiek

0

100

70 Grade of similarity

▶ głowie	2
▶ łowicka	1
▶ człowiek	8
▶ chodzącem	4

☐ continue to build a complex query

Cancel

Back

Next

Finish

Ilustracja 2.23. Wyświetlanie wyników wyszukiwania przeprowadzonego według stopnia podobieństwa

Źródło: opracowanie własne na podstawie programu CATMA

Jeśli chcemy zmienić maksymalną liczbę wyników wyświetlanych w podglądzie, należy wpisać żadaną maksymalną wartość w dolnej części Kreatora zapytań i kliknąć opcję Wyświetl w podglądzie.

Jeśli wyniki zawierają różne typy słów, które spełniają nasze zapytanie, słowa te będą wyświetlane w różnych wierszach wraz z ich odpowiednią częstotliwością. Na dole sekcji podglądu widać całkowitą liczbę typów słów, które pasują do zapytania i całkowitą częstotliwość ich wystąpień w tekście. Po zakończeniu klikamy przycisk *Finish*, aby wyświetlić wyniki kwerendy w Analizatorze.

Uwaga: Niestety w przypadku języka polskiego narzędzie to bywa zawodne i nie należy bardzo dokładnie sprawdzać wyniki wykonanego w ten sposób przeszukiwania.

2.7.3.4. Wyszukiwanie za pomocą tagów

Kolejną opcją dostępną w programie CATMA jest wyszukiwanie według tagów (znaczników). W tym wypadku, aby wykonać tego rodzaju przeszukiwania, należy wybrać opcję *by Tag* i kliknąć przycisk *Next*. Wykonanie tych czynności sprawi, że wyświetli się lista zestawów znaczników zawierających tagi, które zostały już użyte do oznaczania tekstu. Teraz należy kliknąć na małą strzałkę przed zestawem znaczników, aby wyświetlić odpowiednie tagi.

Uwaga: Ten rodzaj kwerendy jest możliwy tylko po uprzednim „otagowaniu” tekstu.

Na przykład: chcemy wyszukać wszystkie części tekstu, które zostały oznaczone tagiem o nazwie „wsparcie”, który jest częścią zestawu znaczników (*Tagset*) „sport niepełnosprawnych”. Aby to uczynić, należy kliknąć małą strzałkę przed Tagsetem „Oznaki emocji” i zaznaczyć tag „wsparcie”. Po wybraniu jednego ze znaczników podgląd bieżących wyników zapytania zostanie wyświetlony w dolnej części Kreatora zapytań.

W sytuacji gdy chcemy zmienić maksymalną liczbę wyników wyświetlanych w podglądzie, możemy wpisać żądane maksimum w dolnej części Kreatora zapytań i kliknąć na przycisk *Show in preview*. Jeśli wyniki kwerendy zawierają różne typy słów, które spełniają warunki zapytania, słowa te będą wyświetlane w różnych wierszach wraz z odpowiednimi częstotliwościami ich występowania. Na dole sekcji podglądu, po lewej stronie widoczna jest całkowita liczba typów słów, które pasują do zapytania, natomiast po prawej stronie podana będzie całkowita częstotliwość ich wystąpień w tekście.

Query

tag="trudnošći %" Execute Query

Query Builder Wordlist 2

Result by Phrase Result by Tag

Tag	Frequency	Visible in Kwic
▶ /trudnošći	2	☐

Total count: 1
Total frequency: 2

☐ flat table

Ilustracja 2.24. Przykład wyszukiwania za pomocą tagów

Źródło: opracowanie własne na podstawie programu CATMA

Podczas wyszukiwania według tagów można również zdecydować o tym, czy wyniki takiego wyszukiwania mają być prezentowane z wykorzystaniem tagów, czy fraz. Aby to zrobić, należy wybierać zakładkę Wynik przez *Result by markup*. Jeśli wyniki zapytania zawierają różne tagi, każdy tag będzie wyświetlany w innym wierszu.

2.7.3.5. Wyszukiwanie według kolokacji

Jeszcze inną opcją jest wyszukiwanie według kolokacji. Oznacza to, że można szukać wystąpień określonego słowa pod warunkiem, że musi ono pojawić się w pobliżu innego określonego przez nas słowa. Aby skorzystać z tej funkcji przeszukiwania, należy wybierać opcję kolokacji (*by collocation*) i kliknąć przycisk *Next*. Następnie trzeba uzupełnić odpowiednie pola, wpisując w pierwszym z nich słowo, pod względem którego chce się przeszukać tekst, pod warunkiem, że będzie ono współwystępowało obok słowa wpisanego w drugim polu. W trzecim polu można wpisać zakres słów, w ramach którego miałyby się znajdować się te dwa słowa.

Na przykład: możemy wyszukać wszystkie wystąpienia słowa „szpitala”, które pojawiają się w ciągu 20 słów, wraz ze słowem „wypadek”. W tym celu wystarczy wpisać, że chce się wyszukać wszystkie wystąpienia „szpitala”, które pojawiają się w pobliżu „usłyszane” i wybierać 20 jako zakres.

Jeśli klikniemy opcję *Show in preview*, składnia aktualnego zapytania i podgląd bieżących wyników kwerendy zostaną wyświetlone w dolnej części Kreatora zapytań.

I podobnie jak miało to miejsce w przypadku poprzedniego typu przeszukiwania, także i tutaj, jeśli chce się zmienić maksymalną liczbę wyników wyświetlanych w podglądzie, należy wpisać żądane maksimum w dolnej części Kreatora zapytań i kliknąć *Show in preview*. Jeśli wyniki zapytania zawierają różne typy słów, które spełniają warunki zapytania, słowa te będą wyświetlane w poszczególnych wierszach, wraz z ich częstotliwością występowania. Na dole po lewej stronie sekcji podglądu widoczna jest całkowita liczba typów słów, które pasują do zapytania, zaś po prawej stronie całkowita częstotliwość wystąpień słowa.

2.7.3.6. Wyszukiwanie według częstotliwości

Niemniej użyteczną opcją jest wyszukiwanie według częstotliwości, które wykonuje się poprzez wybór opcji *by frequency* i kliknięcie przycisku *Next*. Następnie można określić, wpisując w odpowiednie pola liczby określające przedział, w jakim program będzie sprawdzał częstotliwość występowania słów w tekście.

Na przykład: można wyszukiwać każde słowo, które pojawia się w tekście od dziesięciu do dwudziestu razy. W tym celu należy wybrać „between” („pomiędzy”) z rozwijanej listy i wpisać „10” w pierwszym polu, a „20” w drugim.

Query Builder

1. How do you want to search?

2. Search for all occurrences of

wypadek*
that appear near*
szpitala
within a span of*
10

Show in preview with a maximum total frequency of 50

Your search "wypadek*" & "szpitala" 10 will match for example:

Phrase	Frequency
▶ wypadek	1

☐ continue to build a complex query

Cancel Back Next Finish

Ilustracja 2.25. Okno Kreatora zapytań podczas wyszukiwania przez kolokację

Źródło: opracowanie własne na podstawie programu CATMA

Query Builder

1. How do you want to search?

2. How frequent shall it be?

The word shall appear

between

10

*

and

20

*

times.

Show in preview

with a maximum total frequency of

50

Your search

freq = 10-20

will match for example:

Phrase

była

Frequency

10

☐ continue to build a complex query

Cancel

Back

Next

Finish

Ilustracja 2.26. Okno Kreatora zapytań z włączoną opcją przeszukiwania według częstotliwości

Źródło: opracowanie własne na podstawie programu CATMA

Jeśli kliknie się opcję *Show in preview*, składnia aktualnego zapytania i podgląd bieżących wyników kwerendy zostaną wyświetlone w dolnej części Kreatora zapytań. Można także ustalić maksymalną liczbę wyników wyświetlanych w podglądzie, wpisując żądane maksimum w dolnej części Kreatora zapytań i klikając *Show in preview*. Jeśli wyniki kwerendy zawierają różne typy słów, które spełniają zapytanie, słowa te będą wyświetlane w poszczególnych wierszach wraz z odpowiednimi częstotliwościami. Na dole sekcji, po lewej stronie podglądu widoczna będzie całkowita liczba typów słów, które pasują do zapytania, a po prawej stronie całkowita częstotliwość wystąpień w tekście.

2.7.4. Zapytania złożone z zastosowaniem kreatora kwerend

Poza wymienionymi już rodzajami przeszukiwania danych program CATMA pozwala także na bardziej skomplikowane sposoby ich konstruowania. Generalna zasada, o której należy pamiętać, chcąc realizować złożone zapytania, jest następująca. Otóż w momencie, gdy wykonaliśmy już jedno z wyżej wymienionych zapytań podstawowych, zamiast klikać na przycisk *Finish*, należy zaznaczyć pole *continue to build a complex query* znajdujące się w prawym dolnym rogu Kreatora zapytań (*Query Builder*) i w ten sposób kontynuować budowanie złożonych zapytań.

Kolejną czynnością jest wybór sposobu, w jaki możemy rozszerzyć swoje zapytanie. Może to nastąpić poprzez: wprowadzenie większej ilości wyników, zastosowanie dodatkowych kryteriów selekcji (wyłączania zapytań) lub modyfikowanie poprzednio uzyskanych wyników wyszukiwania. Po dokonaniu wyboru którejs z opcji klikamy na przycisk *Next*.

Niezależnie od tego, jaki sposób budowania złożonych zapytań wybierzemy, zawsze należy zdecydować o tym, czy kontynuowanie przeszukiwania będzie dotyczyło słów bądź frazy, czy też będzie się odbywało według stopnia podobieństwa, przez przeszukiwanie tagów bądź kolokacje i częstotliwości występowania danych.

2.7.4.1. Wprowadzenie dodatkowych kryteriów przeszukiwania

Jeśli zdecydujemy o użyciu opcji wprowadzenia dodatkowych kryteriów przeszukiwania, będzie to w praktyce oznaczało połączenie dwóch wybranych zapytań. Na przykład: możemy w ten sposób wyszukać każde pojawienie się słowa „szpital”, a jednocześnie każde słowo, które występuje więcej niż 3 razy w tekście. W tym celu należy wykonać następujące kroki. Po pierwsze wybieramy opcję *by word or phrase*, na początku ustalamy parametry przeszukiwania i klikamy przycisk *Next*. Po drugie wpisujemy słowo „szpital” i zaznaczamy opcję *continue to*

Query Builder

1. How do you want to search?

2. How does your phrase look like?

3. What do you want to follow it in next query?

The first word starts with

szpital

The first word contains

The first word ends with

This first word is adjacent to

Show in preview

with a maximum total frequency of

3

Your search

wild="szpital%"

will match for example:

Phrase	Frequency
▶ szpital	1
▶ szpitala	7

☒ continue to build a complex query

Cancel

Next

Finish

Ilustracja 2.27. Przykład zastosowania zapytania opartego na wprowadzeniu dodatkowych kryteriów przeszukiwania

Źródło: opracowanie własne na podstawie programu CATMA

build complex query, ponownie klikając na *Next*. Po trzecie wybieramy *add more results* i tak jak poprzednio klikamy *Next*. Po czwarte decydujemy o wyborze kryterium *frequency*, co także potwierdzamy, naciskając na *Next*. Wreszcie, po piąte wybieramy opcję *greaterThan* z rozwijanej listy, wpisujemy „3” w odpowiednim polu i klikamy na *Show in preview*. Całą procedurę zamykamy, naciskając na przycisk *Finish*.

2.7.4.2. Zastosowanie dodatkowych kryteriów selekcji

Zastosowanie dodatkowych kryteriów selekcji będzie *de facto* oznaczało wykluczenie niektórych działań, które dotyczą wybranego rodzaju zapytania. Na przykład: możemy wyszukiwać wszystkie słowa, które są w 50% podobne do słowa „niepełnosprawność”, ale tylko wtedy, gdy pojawiają się przynajmniej 3 razy w tekście. W tym celu należy wykonać następujące kroki. Po pierwsze należy wybrać przeszukiwanie według stopnia podobieństwa i potwierdzić wybór przyciskiem *Next*. Po drugie trzeba wpisać słowo „niepełnosprawność” oraz określić skalę podobieństwa na 50 (procent). Po trzecie zaznaczamy opcję *continue to build complex query* i potwierdzamy to przyciskiem *Next*. Po czwarte wybieramy *exclude hits from previous results*, również klikając na *Next*. Po piąte określamy częstotliwość i wybieramy opcję *lessThan* z rozwijanej listy oraz wpisujemy w odpowiednie pole „3”, ponownie potwierdzając nasz wybór przyciskiem *Next*. Wreszcie, po siódme klikamy na *Show in preview* i całą procedurę zamykamy, naciskając na *Finish*.

2.7.4.3. Modyfikowanie poprzednio uzyskanych wyników wyszukiwania

W tym przypadku modyfikowanie zapytania będzie polegać na zawężaniu poprzednio uzyskanych wyników, co może nastąpić, gdy wprowadzone zostaną dodatkowe kryteria przeszukiwania. Na przykład: możemy wyszukać wszystkie wystąpienia słowa „stracić” (i wszystkich jego odmian i form), które są również oznaczone tagiem „trudności”. Aby to zrobić, należy wykonać następujące kroki. Po pierwsze należy wybrać opcję *by word or phrase* i kliknąć na *Next*. Po drugie wpisujemy frazę „trac” (jako element słowa „stracić”) i wybieramy *continue to build a complex query*, co potwierdzamy przyciskiem *Next*. Po trzecie zaznaczamy *refine previous results* i również klikamy na *Next*. Po czwarte wybieramy opcję *by Tag* oraz naciskamy na *Next*. Po piąte klikamy na małą strzałkę przed zbiorem tagów (Tagset), w którym znajduje się znacznik „trudności”. Całą procedurę zamykamy, wybierając *Finish*.

Query Builder

1. How do you want to search?

2. The word is similar to

3. What do you want to do with the next query?4. How do you want to search?

5. How does your phrase look like?

niepełnosprawność

0100

50 Grade of similarity

▶ niepełnosprawnymi	1
▶ ukierunkować	1
▶ niepełnospr	1
▶ niepełnosprawność	1
▶ pełnosprawni	1

☒ continue to build a complex query

Cancel

Back

Next

Finish

Ilustracja 2.28. Przykład zastosowania zapytania opartego na zastosowaniu dodatkowych kryteriów selekcji przeszukiwania danych

Źródło: opracowanie własne na podstawie programu CATMA

Query Builder

+ x

1. How do you want to search?2. How does your phrase look like?3. What do you want to do with the next query?4. How do you want to search?5. Please choose a Tag (only Tags that v

Tagsets	Tag Color
▼ sport osób niepełnosprawnych	
trudności	

Show in preview

with a maximum total frequency of 50

Your searchtag="trudności %"

will match for example:

Phrase	Frequency
▶ ten cewnik straszne też problemy stwarzał, ponieważ tam jakiś chyba pewnie zaka	1
▶ takie, że ja żyję i właśnie tak wygląda piekło, bo nawet nadzieję wtedy straclem, to j	1

☐ continue to build a complex query

Cancel

Back

Next

Finish

Ilustracja 2.29. Przykład zastosowania zapytania opartego na modyfikowaniu poprzednio uzyskanych wyników wyszukiwania

Źródło: opracowanie własne na podstawie programu CATMA

2.7.4.4. Kontrola i modyfikowanie poprawności złożonych zapytań

Wykonując złożone zapytania, trzeba pamiętać, że Kreator zapytań umieszcza specjalne nawiasy w celu określenia kolejności przetwarzania informacji. Należy zwrócić uwagę na to, że składnia zapytania działa jak formuła matematyczna: informacje w nawiasach są przetwarzane jako pierwsze. Może się jednak zdarzyć, że Kreator zapytań nie umieści nawiasów w taki sposób, który nam odpowiada. W takim przypadku należy odręcznie zmienić nawiasy.

Jeśli na przykład uruchomisz Kreator zapytań z opcją *by word or phrase* i wpiszesz „s” w polu *The first word starts with* (pierwsze słowo zaczyna się od), i dodatkowo doprecyzujesz zapytanie przy użyciu tagu „trudności”, wybierzesz *exclude hits from previous results* (wykluczenie trafień z poprzednich wyników) i wykluczysz wszystkie słowa z więcej niż pięcioma wystąpieniami, i dodasz więcej wyników, stosując opcje *by word or phrase* (ponownie używając opcji słowo lub fraza), ale tym razem wpiszesz „t” w polu *The first word starts with* (pierwsze słowo zaczyna się od), wówczas formuła zapytania będzie wyglądała tak jak poniżej:

```
((reg = „\b\Qs\E\S*”) where (tag = „Okrzyk”)) – (freq > 5)),  
(reg = „\b\S*\Qt\E(? = \W)”)
```

Kwerenda wyszuka wszystkie wystąpienia słów, które zaczynają się od litery „s” (reg = „\b\Qs\E\S*”) i pojawiają się w częściach tekstowych oznaczonych jako „niepełnosprawność” (gdzie (tag) = „trudności”), przy czym, wszystkie te słowa, które występują więcej niż pięć razy, są wykluczone (– (freq > 5)), a dodatkowo wyszukiwanie dotyczy wszystkich wystąpień zakończonych literą „t” (reg = „\b\S*\Qt\E(? = \W)”).

Może się tak zdarzyć, że będziemy chcieli wyszukać wszystkie wystąpienia słów, które zaczynają się od litery „s” (reg = „\b\Qs\E\S*”) i pojawiają się w częściach tekstowych oznaczonych jako „niepełnosprawność” (where (tag = „trudności”)), przy czym chcemy wykluczyć nie tylko wszystkie słowa, które występują więcej niż pięć razy, ale także wszystkie wyrazy kończące się na „t”. W takim razie należy ręcznie zmienić nawiasy klamrowe:

```
(reg = „\b\Qs\E\S*” where tag = „Okrzyk”) – (freq > 5,  
reg = „\b\S*\Qt\E(? = \W)”)
```

W przypadku tego zapytania ważne jest, aby wyjaśnić, że pierwsza część („Wyszukuję słowa zaczynające się od litery „s”, które pojawiają się w częściach tekstowych oznaczonych jako „niepełnosprawność”) to jedna jednostka. Dlatego są wokół niej nawiasy klamrowe ((reg = „\b\Qs\E\S*” where tag = „trudności”))

– i chcesz wykluczyć inną jednostkę (słowa, które pojawiają się więcej niż pięć razy i słowa, które kończą się na literę „t”). Dlatego wstawiane są również „(freq > 5, reg = „\b\S*\Qt\E(? = \W)”)” w nawiasach.

2.7.5. Kwerenda wykonywana za pomocą bezpośrednich zapytań

Mimo że większość zapytań daje się wygenerować za pomocą Kreatora zapytań, można również wykonać zapytanie bezpośrednio. Aby to zrobić, należy wpisać zapytanie w polu w lewym górnym rogu Analizatora i kliknąć opcję *Execute Query* (Wykonaj kwerendę).



Ilustracja 2.30. Pole *Execute Query* w oknie Analizatora

Źródło: opracowanie własne na podstawie programu CATMA

Uwaga: Należy pamiętać, aby zawsze sprawdzać zapytania generowane przez Kreatora zapytań. Ma to podwójne znaczenie. Po pierwsze powala sprawdzić ich prawidłowość i ocenić, czy dana formuła spełnia nasze oczekiwania. Po drugie daje to możliwość przyzwyczajania się do specjalnych wyrażeń i składni, które są niezbędne do osiągnięcia określonych rezultatów.

2.7.5.1. Kwerenda oparta na prostych bezpośrednich zapytaniach

Program CATMA umożliwia wykonywanie przeszukiwania danych z zastosowaniem tak zwanych „prostych zapytań”. Zalicza się do nich między innymi: narzędzia służące do wyszukiwania słów i fraz (bądź ich części w zgromadzonych dokumentach), a także narzędzia oparte na systemie symboli zastępujących fragmenty wyrazów oraz położenie wyrazów w ramach danego tekstu.

2.7.5.1.1. Zapytania ze słowami lub frazami

Najbardziej podstawowy typ zapytania, który można stworzyć, składa się tylko z jednego słowa. Przy czym należy pamiętać, że zawsze takie słowo musi się znaleźć w cudzysłowie.

Na przykład:

„niepełnosprawność”

Rezultatem przeprowadzonej na podstawie tej formuły kwerendy byłoby wyszukanie wszystkich wystąpień słowa „niepełnosprawność” w danym tekście.

Budowanie zapytań ze zwrotami jest równie proste. Wystarczy wpisać frazę i umieścić ją w całości w cudzysłowie.

Na przykład:

„Niepełnosprawność stała się po prostu częścią mnie”

W ten sposób wyświetlą się wszystkie wystąpienia wyszukiwanej frazy.

Uwaga: Ponieważ odwrócone przecinki mają specjalne znaczenie dla analityka zapytań, nie można ich używać za wyjątkiem ich faktycznego przeznaczenia. Dlatego, aby uniknąć nieporozumień i błędów w wyszukiwaniu danych, należy stosować lewe ukośniki, które niwelują ich specjalne zastosowania. Na przykład: „\” Kuba! \” „Wrzasnąłem”

2.7.5.1.2. Zapytania z wyrażeń regularnych

Tak zwane wyrażenia regularne umożliwiają wyszukiwanie słów oraz ich części. Kiedy chce się utworzyć zapytanie z wyrażeniem regularnym, zapytanie należy rozpocząć od formuły:

reg =

i wstawić następujące zapytanie w cudzysłowie.

W poniższej tabeli znajduje się lista możliwych operatorów dla zapytań z wyrażeń regularnych wraz z opisem ich funkcji:

Operator	Opis zastosowania
.	Kropka reprezentuje dowolny pojedynczy znak Na przykład: jeśli nasze zapytanie będzie następującej postaci reg = „k..t” , oznacza to, że szukamy dowolnej sekwencji czterech znaków, gdzie po znaku „k” następują dwa dowolne znaki i litera „t”. Należy pamiętać, że kropka może również reprezentować puste pola (spacje). Dlatego w wyniku wyszukiwania możemy uzyskać „kret”, ale także „ki t” (jako część wyrażenia „ciężki tornister”).

Operator	Opis zastosowania
*	Gwiazdka reprezentuje brak, jedno lub więcej następujących po sobie wystąpień poprzedzającego znaku Na przykład: jeśli nasze zapytanie będzie następującej postaci reg = „bis*” , oznacza to, że szukamy dowolnej sekwencji liter „bi”, a następnie litery „s” niewystępującej lub występującej co najmniej jeden raz. Dlatego w wyniku wyszukiwania możemy uzyskać „bi” (jako część słowa „bis”) czy „bis” (jako część słowa „biskopt”).
?	Znak zapytania reprezentuje brak lub jedno kolejne wystąpienie poprzedniego znaku Na przykład: jeśli wpisujemy formułę reg = „dis?” , możemy uzyskać „di” i „dis” wśród wyników.
+	plus podobnie jak znak zapytania reprezentuje brak lub jedno, lub więcej wystąpień poprzedniego znaku Na przykład: jeśli wpisujemy reg = „dis +” , jako wynik możemy uzyskać „dis” i „diss”.
\b	znak \b reprezentuje znacznik końca słowa Na przykład: jeśli wpisujemy formułę reg = „\bk..t\b” , to wśród wyników możemy uzyskać „kret”, ale nie „ki t”.
[]	nawiasy kwadratowe reprezentują zestaw określonego rodzaju znaków (np. samogłosek) Na przykład, reg = „b [aeiou]” wyświetli każde wystąpienie litery „b”, gdzie po „b” następuje samogłoska, np. „Be” lub „ba” (jako część słowa „balon”).
^	znak ^ służy do zanegowania dowolnej klasy znaków Na przykład reg = „b ^ aeiou” wyświetli każde wystąpienie litery „b”, gdzie po „b” nie następuje samogłoska, np. „By” lub „br” (jako część słowa „brzoza”).
-	minus służy do określenia zakresu w klasie znaków Na przykład, reg = „b [a-m]” wyświetli każde wystąpienie litery „b”, gdzie po „b” następuje litera leżąca między literami „a” i „m” w alfabecie, w tym „a” i „m”, jak „be” lub „bl” (jako część słowa „blenda”).

2.7.5.1.3. Zapytania z zastosowaniem symboli

Niektóre funkcje zapytań z użyciem wyrażeń regularnych (*Queries with Regular Expressions*) można także uzyskać za pomocą zapytań z tak zwanym wieloznacznym symbolem (*Wildcard Queries*). Pozwalają one wyszukiwać słowo, wpisując tylko jego części. Przy czym należy pamiętać, że zawsze, gdy chce się utworzyć zapytanie z symboli wieloznacznych, każde zapytanie trzeba rozpocząć od:

wild =

i wstawić zapytanie w cudzysłowie.

Poniżej znajduje się lista operatorów dla zapytań typu *wildcard* (wieloznacznym symbolem) z opisem ich odpowiednich funkcji:

Operator	Funkcja
%	% procent: reprezentuje nieznaną część słowa Na przykład: jeśli wpisujemy formułę wild = „% ak”, otrzymamy wszystkie słowa, które kończą się na „tak”, np. „tartak”, „barak”, „hamak”. Uwaga: Operator ten może być również używany w Zapytaniach z tagami.
_	_ podkreślenie: reprezentuje dowolny pojedynczy znak w słowie Na przykład: jeśli wpisujemy wild = „_ ak”, otrzymamy dowolne trzyliterowe słowo, które kończy się na „ak”.

Uwaga: Jeśli będziemy chcieli wyszukać znak, który może być również używany jako operator w kwerendzie, wówczas niezbędne okaże się umieszczenie przed nim lewego ukośnika.

Na przykład: jeśli chcemy wyszukać dowolne słowo zaczynające się od „aim_”, należy wpisać wild = „aim _ %”.

2.7.5.1.4. Zapytania oparte na podobieństwie danych

Program umożliwia również przeszukiwanie pod względem podobieństwa danych słów do siebie. Jeśli chce się to zrobić, należy wówczas wstawić

simil =

przed zapytaniem, które należy umieścić w cudzysłowie.

Za twórcami programu można dodać, że wykorzystuje on algorytm rozpoznawania wzorców Ratcliff/Obershelp w celu określenia podobieństwa. Przy czym omówienie algorytmu wykracza poza zakres rozwijanych w tej książce zagadnień. Osoby zainteresowane mogą się jednak odwołać do informacji zawartych na stronie: <http://www.drdoobs.com/database/pattern-matching-the-gestalt-approach/184407970> ? pgno = 5) Oto przykład działania zapytań o podobieństwa:

simil = „man” 70%

Ta kwerenda wybiera dowolne słowo, które jest w 70% podobne do słowa „man”. Należy przy tym zwrócić uwagę, że symbol procentu jest opcjonalny.

2.7.5.1.5. Zapytania z tagami

Kwerendy danych mogą być także prowadzone w oparciu o istniejące tagi. Jeśli bowiem otagowaliśmy tekst, nad którym pracujemy, będzie można użyć tych

metainformacji w procedurze przeszukiwania. Aby wykorzystać tę możliwość, należy rozpocząć formułę wyszukiwania od:

tag =
a samo zapytanie wstawić w cudzysłowie.

Oto przykład:

tag = „trudności”

Przeprowadzona na tej podstawie kwerenda spowoduje, że wyszukane zostaną dowolne słowa lub frazy oznaczone tagiem „trudności”.

Jeśli natomiast chcemy wyszukać określony znacznik stanowiący podkategorię dla danego tagu (subtag), wówczas należy wpisać całą ścieżkę, która prowadzi dożądanego podtagu.

Na przykład:

tag = „trudności/problemy emocjonalne”

To zapytanie wyświetli listę wszystkich wystąpień subtagu „problemy emocjonalne” należącego do znacznika „trudności”.

Możliwe jest również użycie operatora procentu z zapytań typu wildcard podczas wyszukiwania opartego na tagach.

Na przykład:

tag = „%”

Taka kwerenda wyświetli w tekście każde wystąpienie znacznika o nazwie „problemy emocjonalne” niezależnie od tego, do którego znacznika należy podtag „pojedyncze słowo”, ponieważ może istnieć więcej niż jeden subtag o tej nazwie.

2.7.5.1.6. Zapytania oparte na kolokacji danych

Zapytania z kolokacją umożliwiają wyszukiwanie słów, które pojawiają się w pobliżu innych słów. W ten sposób można na przykład dowiedzieć się, gdzie słowo „niepełnosprawność” pojawia się w pobliżu słowa „sport”. Można to łatwo zrobić za pomocą operatora &.

Operator	Opis
&	&: służy do definiowania zapytań kolokacyjnych. Na przykład: formuła „ niepełnosprawność i „ sport ” 20 spowoduje, że wyświetlone zostaną wszystkie wystąpienia słowa „niepełnosprawność”, gdzie „sport” pojawi się w ramach z góry określonej liczby słów po obu stronach wyrazu „niepełnosprawność”. W tym przypadku „20” oznacza wielkość zakresu dla kolokacji. Jeśli pominiesz rozmiar zakresu, domyślną wartością jest pięć słów po obu stronach.

Uwaga: Należy pamiętać, że kwerenda oparta na kolokacjach wyszukuje słowa według reguły „od prawej do lewej”, co oznacza, że słowo znajdujące się skrajnie po lewej stronie będzie zawsze tym, które jest wybrane. Jeśli jednak chcemy przeprowadzić kwerendę opartą na kolokacjach z więcej niż dwoma komponentami, wówczas należy umieścić nawiasy, aby określić, która część kwerendy ma zostać wykonana jako pierwsza.

2.7.5.1.7. Zapytania o częstotliwość występowania danych

Program CATMA umożliwia również tworzenie zapytań, które pozwalają na wyszukiwanie wszystkich wyrazów występujących w tekście z pewną częstotliwością. Aby wykorzystać tę możliwość, należy rozpocząć formułę wyszukiwania od:
freq =

Oto kilka przykładów możliwości zapytań związanych z częstotliwością występowania danych:

Operator	Opis
=	znak równości pomaga wyszukać wyrazy pojawiające się w tekście z dokładnie taką częstotliwością, jaką określimy Na przykład: freq = 5 wyświetli wszystkie słowa występujące pięciokrotnie w tekście.
<	znak mniejszy niż pomaga wyszukiwać słowo występujące w tekście rzadziej niż wybrana częstotliwość Na przykład: freq < 100 wyświetli wszystkie słowa, które występują mniej niż 100 razy w tekście Uwaga: ten operator może być również łączony ze znakiem równości Na przykład: freq <= 100 wyświetli listę wszystkich słów występujących 100 lub mniej razy w danym tekście
>	znak większy niż pomaga wyszukiwać słowa, które występują w tekście częściej niż wybrana częstotliwość Na przykład: freq > 50 wyświetli wszystkie słowa, które występują więcej niż 50 razy w tekście Uwaga: ten operator może być również łączony ze znakiem równości Na przykład: freq >= 100 wyświetli wszystkie słowa, które występują 50 lub więcej razy w tekście
–	minus pomaga wyszukać zakres częstotliwości Na przykład: freq = 5 – 10 wyświetli wszystkie słowa, które pojawią się w tekście od pięciu do dziesięciu razy.

2.7.5.2. Kwerenda oparta na złożonych zapytaniach bezpośrednich

Złożone kwerendy pozwalają łączyć poznane dotychczas operatory i tworzyć na ich podstawie zaawansowane formuły przeszukiwania. W kolejnych podpunktach w zwięzły sposób opisano najważniejsze mechanizmy ich tworzenia oraz podano przykłady bezpośredniego zastosowania.

Uwaga: W przypadku zapytania złożonego, które składa się z więcej niż dwóch komponentów, konieczne jest umieszczenie nawiasów w celu określenia, która część zapytania ma być wykonywana jako pierwsza.

2.7.5.2.1. Łączenie zapytań

Jedną z możliwości, jakie daje program, jest łączenie wyników wielu podstawowych zapytań. Zasada działania jest tutaj bardzo prosta, bowiem łączenie zapytań można osiągnąć za pomocą zwykłego przecinka (,).

Jednym z przykładów może być: „głośniej”, „słyszałem”

Wynikiem takiego zapytania będzie wyświetlenie wszystkich wyrazów „głośniej” i wszystkie wystąpienia „słyszałem”.

Bardziej złożonym przykładem byłoby:

wild = „% ki”, freq > 20

Wynikiem takiego zapytania będzie wyświetlenie wszystkich słów kończących się literami „ki”, a także każdego słowa, które pojawi się w tekście ponad dwadzieścia razy.

2.7.5.2.2. Zapytania z wykluczeniami

Zapytania z tak zwanymi wkluczeniami są pomocne, gdy wyniki przeszukiwań zawierają dużo informacji, a my jesteśmy zainteresowani tylko niektórymi z nich. Ich zmniejszenie i doprecyzowanie można wówczas osiągnąć za pomocą myślnika (–) jako operatora.

Na przykład chcemy sprawdzić występowanie słowa „żał”, ale tylko w sytuacji, w której nie jest ono oznaczone tagiem „emocje”. Takie zapytanie miałoby następującą postać:

„żał” – tag = „emocje”

W ten sposób wyświetlone zostaną wszystkie wystąpienia słowa „żał”, które nie są oznaczone tagiem „emocje”.

2.7.5.2.3. Zapytania z przyległością

Zapytania z przyległością umożliwiają wyszukiwanie określonych wyników, które będą wyświetlane tylko wtedy, gdy bezpośrednio po nich następują wyniki innego zapytania. Kwerendy te są budowane za pomocą średnika (;) używanego jako operatora.

Na przykład:

„niepełnosprawność”; wild = „i%”

Rezultatem tak przeprowadzonej kwerendy jest wyświetlenie wszystkich wystąpień słowa „niepełnosprawność”, po których bezpośrednio następuje słowo zaczynające się na „i” (np. gdy chcemy wyszukać frazy „niepełnosprawność intelektualna”).

Uwaga: W obecnej wersji program nie jest jeszcze w stanie wybrać elementu innego niż pierwszy element w sekwencji.

2.7.5.2.4. Zapytania poszerzone

Poza już wymienionymi istnieje również kilka innych możliwości dopracowania zapytania, np. poprzez użycie formuły podobieństwa według wybranego tagu bądź ze względu na częstotliwość występowania danych. Aby zbudować zapytania z tego rodzaju ulepszeniami, należy użyć skrótu „**where**” jako operatora.

Przykładem takiego zapytania jest:

wild = „a%” where simil = „człowiek” 50%

Rezultatem tak przeprowadzonej kwerendy jest zestawienie zawierające wszystkie wyrazy zaczynające się od litery „a”, które są w 50% podobne do wyrazu „człowiek”.

Innym przykładem może być zapytanie typu:

wild = „\ te%” where freq = 100

Rezultatem kwerendy opartej na powyższym przykładzie jest wyszukiwanie wszystkich słów zaczynających się od „te”, które występują dokładnie 100 razy w tekście.

Kolejną możliwością jest wyszukiwanie za pomocą tagu, które może przybierać postać:

„żał”, where tag = „emocje”

Takie przeszukiwanie będzie pokazywało wszystkie wystąpienia słowa „żał”, które są oznaczone tagiem o nazwie „emocje”.

Warto przy tym pamiętać, że w celu sprecyzowania zapytania za pomocą tagów można również zastosować trzy warianty przeszukiwania: wyszukiwanie „dokładnie” tego, czego chcemy, określające wartości „graniczne” oraz oparte na systemie „nakładania się” danych.

Przykładem zapytania opartego na zasadzie dokładnego przeszukiwania może być formuła:

„żał”, where tag = „emocje” exact

Rezultatem tak przeprowadzonej kwerendy będzie wyszukanie wszystkich wystąpień słowa „żał”, gdzie dokładnie to słowo jest oznaczone tagiem „emocje”.

Z kolei przykładem dopasowania „granicznego” może być:

„żał”, where tag = „emocje” boundary

W tym wypadku wynikiem wyszukiwania będzie wyświetlenie się słowa „żał”, ale tylko wówczas, gdy będą one w całości uwzględnione w tagu o nazwie „emocje”. Będzie tak w przypadku, gdy „żał” i „emocje” są dokładnie dopasowane, ale także, gdy całe słowo „żał” jest oznaczone tagiem „emocje”, który swoim zakresem obejmuje także inne słowa oprócz wyrazu „żał”.

Natomiast przeszukiwanie związane z „nakładaniem się” może wyglądać w następujący sposób:

„żał”, where tag = „emocje” overlap

Kwerenda oparta na tej formule wyszuka wszystkie fragmenty tekstu, w których występuje słowo „żał”, a zarazem są to części oznaczone tagiem „emocje”, które przynajmniej w jakimś zakresie pokrywają się ze sobą. Przy czym słowo i tag nie muszą się dokładnie ze sobą pokrywać ani też wszystkie wystąpienia w tekście słowa „żał” nie muszą być oznaczone tagiem „emocje”.

Można również tworzyć złożone zapytania, używając do tego celu przecinka (,) jako operatora logicznego „i” oraz pionowej kreski (|) jako „lub” operatora „lub”.

Przykładem zastosowania przecinka jako operatora „i” może być formuła:

„żał”, where tag = „emocje”, tag = „negacja”

Rezultatem tak przeprowadzonego zapytania jest wyświetlenie wszystkich wystąpień słowa „żał”, które są jednocześnie oznaczone tagiem „emocje”, jak i tagiem „negacja”.

Z kolei przykładem dla operatora zastosowania pionowej kreski jako operatora „lub” jest:

„żał”, where tag = „emocje” | tag = „niepokój”

Rezultatem tego zapytania będzie wyświetlenie się wszystkich wystąpień słowa „żał”, które są zarazem oznaczone jako „emocje” lub „niepokój”.

Warto również zwrócić uwagę na to, że w przypadku wielokrotnego zastosowania przecinka ma on pierwszeństwo przed operatorem pionowej kreski. Dlatego należy użyć nawiasów, aby wyniki przeszukiwania były zgodne z naszymi oczekiwaniami, w sytuacji gdy pierwszeństwo operatora zmieniłoby zamierzone znaczenie.

2.7.5.2.5. Podsumowanie zagadnień związanych z operatorami służącymi do tworzenia złożonych zapytań

Operator	Opis
,	przecinek łączy wynik, za wyjątkiem, gdy jest stosowany w ramach zapytania złożonego. Wówczas jego działanie jest tożsame z operatorem logicznym „i”
	pionowa kreska jest używana w ramach złożonych zapytań jako operator logiczny „lub”
–	łącznik służy do definiowania wykluczeń (Uwaga: w przypadku wykorzystania łącznika w zapytaniu z częstotliwością może być użyty do zdefiniowania zakresu i częstotliwości.)
;	średnik służy do definiowania sąsiedztwa (przestrzennej lokalizacji wyrazów w tekście)

2.8. Słowa kluczowe w kontekście

Z każdego wykazu wygenerowanego w Analizatorze takiego jak listy słów (*Wordlists*) lub wyniki wyszukiwania (*Query results*) można wybrać pojedyncze słowa lub frazy i wyświetlić je w ramach danego tekstu. Innymi słowy program umożliwia powrót do danych poprzez pokazywanie słów oraz fraz w kontekście ich występowania. Aby to zrobić, należy zaznaczyć pola pod słowami z listy, które chce się umieścić w ich kontekście w tak zwanej kolumnie *Kwic* (skrót od kluczowego słowa w kontekście). Można także zaznaczyć lub odznaczyć wszystkie słowa z danej listy jednocześnie za pomocą przycisków w lewym dolnym rogu okna Analizatora. Słowa wybrane dla *Kwic* zostaną wyświetlone w prawym dolnym rogu okna wraz z dziesięcioma słowami, które je otaczają.

Jeśli chcemy zobaczyć słowo w szerszym kontekście niż ten wyświetlany w *Kwic*, należy dwukrotnie kliknąć na przycisk *Select all for Kwic* znajdujący się w prawym dolnym rogu okna Analizatora, co spowoduje otwarcie Taggera i przejście do odpowiedniej pozycji w tekście.

Widok *Kwic* może także służyć do przeprowadzania operacji tagowania słów. Aby to zrobić, należy wybierać dane słowo z listy *Kwic* za pomocą myszy, a następnie otworzyć Menedżer tagów i przeciągnąć, a następnie upuścić żądany tag

Query

freq>0

Execute Query

+ New Query

Documents and annotations constraining this search

► W 2

Query Builder

Wordlist

Result by Phrase

Result by Tag

Phrase	Frequency	Visible in Kwic
► człowiek	8	☒
► widzę	8	☐
► nadzieję	8	☐
► ta	8	☐
► były	7	☐
Total count: 1.888		
Total frequency: 9.764		

✖

🔍

📄

🔍

📄

🔍

📄

Select all

Deselect all

Select all for Kwic

Deselect all for Kwic

Document/Annotations	Left Context	Keyword	Right Context
W 2	kosztem czegoś. Całe życie	człowiek	się uczy, no.
W 2	zawsze jakąś granicę. No	człowiek	posiada jakąś granicę i stwierdziłem
W 2	kiedyś, że był taki	człowiek	, który zrobił takie rzeczy
W 2	, na to, jaki	człowiek	jest i tak dalej.
W 2	się. A teraz trochę	człowiek	ma też inne wymaganie od
W 2	. Tak się angażuje ten	człowiek	, nie? Bardzo się
W 2	że tak się angażuje ten	człowiek	. naprawdę. Słucham.

1

30

5 token(s) context

Annotate selected results

Select all

Ilustracja 2.31. Okno programu z listą słów kluczowych, które można wyświetlić w kontekście ich występowania

Źródło: opracowanie własne na podstawie programu CATMA

Query

freq>0

Query Builder **Wordlist** ?

Result by Phrase Result by Tag

Phrase	Frequency	Visible in Kwic
▶ człowiek	8	☒
▶ widzę	8	☒
▶ nadzieję	8	☒
▶ ta	8	☒
▶ były	7	☒
Total count: 1 888		

Tools: [Copy] [Paste] [Print] [Select all] [Deselect all]

Select all for kwic Deselect all for kwic

Documents and annotations constraining this search

▶ W 2

Execute Query

+ New Query

Document/Annotations	Left Context	Keyword	Right Context
W 2	kosztom czegoś. Całe życie	człowiek	się uczy, no.
W 2	zawsze jakaś granica. No	człowiek	posiada jakiejś granice i stwierdza:
W 2	kiedyś, że był taki	człowiek	, który zrobił takie rzeczy
W 2	, na to, jaki	człowiek	jest i tak dalej,
W 2	się. A teraz trochę	człowiek	ma też inne wymagania od
W 2	. Tak się angażuje ten	człowiek	, nie? Bardzo się
W 2	że tak się annozuje ten	człowiek	nabrawnie. Słucham.

Tools: [Copy] [Paste] [Print] [Select all] [Annotate selected results] ?

1 30
5 token(s) context

Ilustracja 2.32. Proces tagowania słów wyświetlanych w kontekście ich występowania

Źródło: opracowanie własne na podstawie programu CATMA

do okna, w którym wyświetlane będzie dane słowo w kontekście (jeśli nie korzystaliśmy jeszcze z Menedżera tagów podczas bieżącej sesji roboczej, musisz otworzyć go za pomocą przycisku Otwórz Bibliotekę znaczników w Menedżerze repozytorium). Operację potwierdzamy w wyskakującym okienku pop-up.

Program CATMA pozwala także na otagowanie jednocześnie kilku wystąpień danego słowa w kontekście. Aby wybrać więcej niż jedno wystąpienie danego słowa, należy przytrzymać klawisz Control na klawiaturze. Jeśli zaś chcemy wybrać wszystkie wystąpienia z programu *Kwic* do tagowania, powinniśmy wybierać pierwsze wystąpienie z listy *Kwic*, a następnie przytrzymać Shift i przewinąć listę do samego końca. W ten sposób wszystkie wystąpienia w *Kwic* powinny zostać wybrane. Teraz wystarczy już tylko przeciągnąć i upuścić żądany Tag z Menedżera tagów do *Kwic* i zatwierdzić całą operację.

2.9. Analiza Korpusu

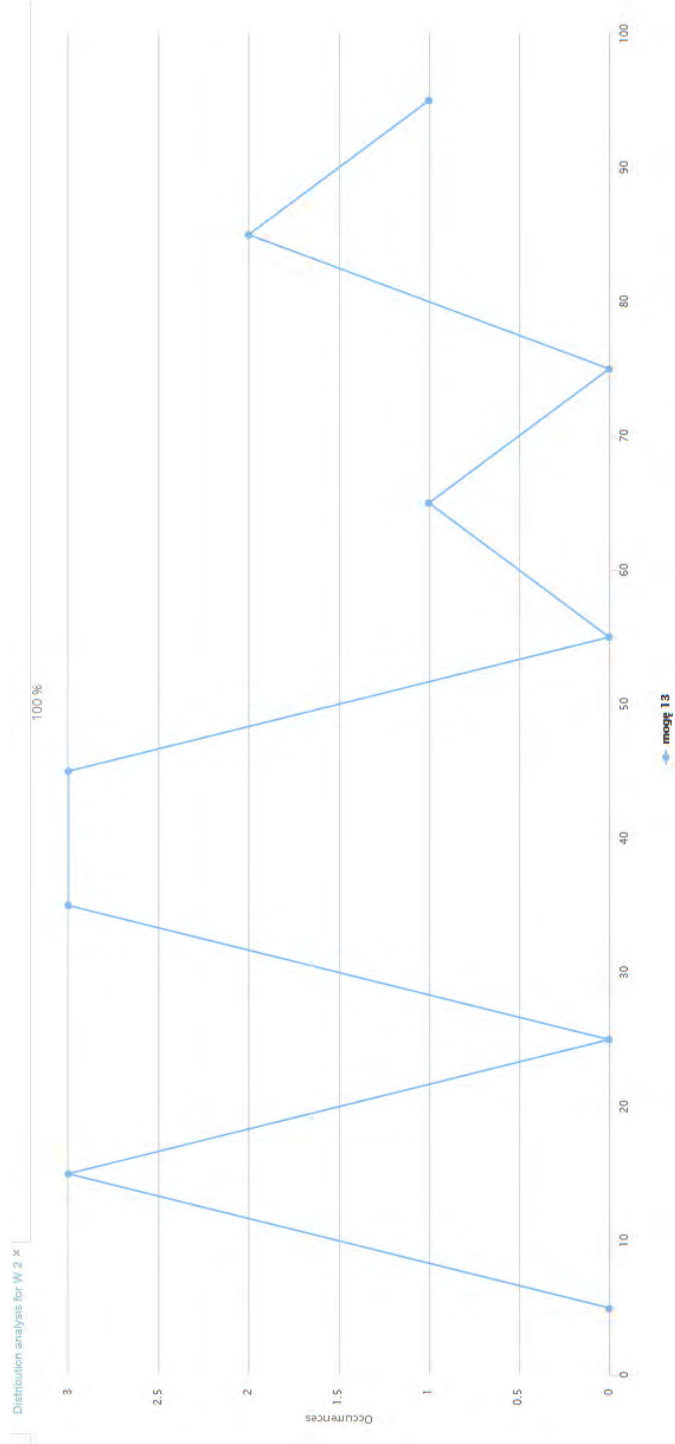
Każdą operację związaną z analizą tekstu, którą można zastosować do pojedynczego dokumentu, można również zastosować do całego korpusu. W tym celu należy wybrać korpus w Menedżerze repozytorium, kliknąć przycisk *More actions...* (Więcej akcji...) poniżej sekcji korpusów oraz wybrać opcję *Analyze Corpus* (Analizuj korpus). Po wyświetleniu listy wyników (Wordlist) lub zapytań (Query) należy kliknąć małą strzałkę przed daną pozycją, aby wyświetlić istniejące zależności między różnymi dokumentami analizowanego korpusu. Analizowane dokumenty będą wyświetlane w prawym górnym rogu okna analizatora.

2.10. Wizualizacja danych

W programie obecny jest również *Visualizer*, który oferuje możliwość wyświetlania rozkładu słów z listy lub wyników zapytań. Samo korzystanie z tego narzędzia jest bardzo proste. Aby otworzyć wizualizator, należy najpierw utworzyć Analityzer i utworzyć listę słów lub uruchomić kwerendę. Przy czym istnieje kilka możliwości tego, jak dane są wyświetlane w ramach wizualizatora, co zostało w skrócie opisane w poniższych podpunktach.

2.10.1. Wizualizacja za pomocą prostego wykresu

Przede wszystkim wizualizacja może polegać na utworzeniu prostego wykresu odwzorowującego rozkład wybranych wyników w tekście. Aby to uczynić, należy wybrać odpowiednią pozycję z dolnej lewej sekcji analizatora, a następnie kliknąć



Ilustracja 2.33. Widok wykresu obrazującego rozkład wybranych wyników w tekście

Źródło: opracowanie własne na podstawie programu CATMA

mały symbol wykresu w lewym dolnym rogu okna analizatora. W ten sposób uzyskamy wykres przedstawiający rozkład wybranych wyników w tekście. Wykres pokazuje liczbę występowania danego słowa w tekstach (oś y) oraz zakres jego występowania w ramach sekcji tekstu (oś x).

Należy przy tym pamiętać, że tekst jest podzielony na dziesięć równych części – jedna sekcja tekstu składa się z dziesięciu procent tekstu. Jeśli przesuniemy kursor nad jedną z kropek na wykresie, wyświetlona zostanie dokładna liczba wystąpień słowa danego typu i sekcja tekstowa. Natomiast dwukrotne kliknięcie kropki spowoduje przejście bezpośrednio do odpowiedniej części tekstowej w oknie Taggera.

2.10.2. Łączenie wyników w jeden wykres

Za pomocą narzędzia wizualizatora można także łączyć różne słowa lub wyniki zapytania w jednym wykresie. Aby wybrać więcej niż jedną pozycję z dolnej lewej części okna Analizatora, należy przytrzymać klawisz Ctrl na klawiaturze. Jeśli zaś chcemy wybrać wszystkie wyniki, trzeba przytrzymać klawisz Shift i przewinąć listę w dół do ostatniego wpisu. Tym sposobem wszystkie wyniki powinny zostać wybrane. Następnie klikamy mały symbol wykresu w lewym dolnym rogu analizatora. Wybrane wpisy będą wyświetlane na jednym wykresie w wizualizatorze. Listę wpisów można znaleźć na przycisku pod wykresem.

2.10.3. Wizualizacja za pomocą wielu wykresów

Możliwe jest również wyświetlanie różnych wykresów w jednym oknie. Aby to uczynić, należy wygenerować wykres rozkładu, jak to zostało opisane powyżej. Następnie trzeba powrócić do analizatora i wygenerować kolejny wykres. Oba wykresy będą teraz wyświetlane na tym samym oknie. Wykresy można rozróżnić na podstawie ich kolorów i symboli kropek.

Poniżej okna znajduje się lista wykresów i wyniki, które one reprezentują. Klikając symbol wykresu, można go ukryć. Jest to przydatne, zwłaszcza gdy wyświetla się bardzo wiele wykresów naraz.

2.11. Uwagi końcowe

CATMA jest bezpłatnym programem dostępnym w najnowszej wersji on-line. Został umieszczony na serwerze obsługiwany przez Uniwersytet w Hamburgu w Niemczech, zaś serwer kopii zapasowych znajduje się na Uniwersytecie w Heidelbergu oraz dwóch dodatkowych serwerach dzierżawionych od komercyjnych dostawców.

Do niewątpliwych zalet programu należy to, że jest to narzędzie stosunkowo uniwersalne i przydatne dla wszystkich tych, którzy decydują się na pracę z tekstem. CATMA to program, który pozwala na zestawianie dokumentów w korpusie, tworząc ich kolekcję. Co istotne, informacje zawarte w programie mogą być współdzielone z innymi osobami (na zasadzie wyłącznie ich odczytu lub możliwości pełnej edycji). Dla codziennego użytkowania ważne jest także to, że dane mogą być zapisywane w formacie XML zgodnym z TEI, a zatem mogą być łatwo eksportowane do innych aplikacji. Jednocześnie wszystkie dane zamieszczane na serwerze są chronione, a twórcy programu zapewniają, że dokładają najwyższych starań, aby zapewnić prywatność danych i ochronę IP. Na podkreślenie zasługują również takie właściwości programu, jak: obsługa tekstów w prawie każdym języku, pełna integracja funkcji do opisu i analizy danych w przeglądarce internetowej, współpraca za pośrednictwem sieci internetowej umożliwiającej łatwą wymianę dokumentów, adnotacji i znaczników, dowolnie definiowalne lub predefiniowane znaczniki czy wbudowana wizualizacja wyników wyszukiwania i analiz, a także analiza skomplikowanych korpusów tekstowych w jednym kroku. Program daje też możliwość pracy pojedynczego badacza lub współpracy w czasie rzeczywistym w ramach zespołu.

Pewnym zarzutem wobec programu może być w niektórych momentach uciążliwy system tworzenia zapytań, który wymaga wprawy użytkownika. Niemniej jednak w miarę korzystania z programu i nabywania doświadczenia w zakresie formułowania zapytań problem ten staje się coraz mniej kłopotliwy i przestaje nastroczać początkowo odczuwanych trudności.

CATMA (*Computer Assisted Textual Markup and Analysis*) jest programem powstałym w wyniku projektu przez Uniwersytet w Hamburgu, Wydział Języka, Literatury i Mediów. Szefem projektu jest Prof. Dr Jan Christoph Meister, a jego koordynatorem Jan Horstmann, M.A.

Zarówno sam program, jak i dodatkowe informacje o nim są dostępne na stronie:

<http://catma.de/>

Informacja o cytowaniu: Meister J.C., Petris M., Gius E.: CATMA 5.0 (2016) [software for text annotation and analysis]

3. CAT – internetowy zestaw narzędzi do analizy danych jakościowych

CAT, a więc Coding Analysis Toolkit to internetowy zestaw narzędzi CAQDAS. Jest to darmowe oprogramowanie o otwartym kodzie źródłowym, opracowane przez Program Jakościowych Analiz Danych Uniwersytetu w Pittsburghu.

CAT jest w stanie importować dane Atlas.ti, ale ma również wewnętrzny moduł kodujący. Ponadto program umożliwia zarządzanie danymi, kodami, a także analizę materiałów tekstowych. Posiada również rozbudowany system raportowania, dzięki czemu możliwe jest generowanie informacji będących rezultatem prowadzonej pracy analitycznej.

Ciekawostką jest, że program został zaprojektowany do używania klawiszy i automatyzacji w przeciwieństwie do kliknięć myszką, aby przyspieszyć zadania CAQDAS.

3.1. Pierwsze kroki – rejestracja użytkownika i system logowania

Ponieważ CAT to program, z którego można korzystać on-line, wymaga on każdorazowo rejestracji nowego użytkownika. Czynność tę wykonuje się na stronie głównej programu, gdzie w jej górnym obszarze znajduje się link do odpowiedniego panelu (*register for a free account*).

Klikając na wskazany link, przechodzimy przez kolejne etapy procedury rejestracji, wpisując kolejno: unikalną nazwę i hasło użytkownika, wymagane do rejestracji informacje (imię i nazwisko oraz e-mail), a następnie potwierdzamy chęć założenia konta, klikając na odpowiedni link. W ten sposób otrzymujemy automatyczną wiadomość e-mail, na podany uprzednio adres skrzynki elektronicznej, w której znajduje się link do autoryzacji konta (po kliknięciu na odpowiednie oświadczenia o prywatności i bezpieczeństwie danych). Po tych czynnościach w ciągu kilku minut powinniśmy otrzymać kolejną wiadomość e-mail, w której autoryzacja nazwy użytkownika oraz hasła do konta zostaną potwierdzone.

Uwaga: Trzeba pamiętać, że przed pierwszym uruchomieniem programu należy określić i utworzyć subkonto użytkownika. Przy czym jest to wymagane wówczas, gdy zamierzamy podejmować działania angażujące więcej niż jedną osobę. Jeśli zaś będziemy jedynym użytkownikiem, to taka operacja nie jest już konieczna.


Coding Analysis Toolkit

[Lubię to! 54](#)

[register for a free account](#) | [forgot password?](#)

[Home](#) | [About CAT](#) | [DiscoverText](#) | [Terms of Service](#) | [Privacy Statement](#) | [CAT Help Wiki](#) | [Contact Us](#)

[What CAT Does](#)
[CAT Features](#)
[Praise for CAT](#)

Welcome to the Coding Analysis Toolkit (CAT)

CAT is a free service originally by the Qualitative Data Analysis Program (QDAP), and was hosted by the University Center for Social and Urban Research, at the University of Pittsburgh, and QDAP-UMass, in the College of Social and Behavioral Sciences, at the University of Massachusetts Amherst. CAT was the 2008 winner of the "Best Research Software" award from the organized section on Information Technology & Politics in the American Political Science Association.

For the CAT Quick Start Guide, you can view the PDF file here:
[CAT Quickstart Guide](#)

To view a tutorial on using CAT, click here:
[CAT Tutorial - February 23, 2009](#)

May 5, 2010 - CAT is now an open source project! You can host your own version of CAT from the project source code at:
<http://sourceforge.net/projects/catoolkit/>

CAT Statistics

There are currently **15,225** primary CAT accounts and **1,604** sub-accounts. CAT users have uploaded **9,701** coded datasets and **17,051** raw datasets. They have coded a total of **2,300,714** items and adjudicators have made **209,313** validation choices in CAT.



If you like CAT, you'll love DiscoverText. DiscoverText is a cloud-based, collaborative text analytics solution. Generate valuable insights about customers, products, employees, news, citizens, and more. [Sign up for a 30 day free trial.](#)

© 2007-2017 - CAT is maintained by [Texifter, LLC](#) and powered by [Microsoft ASP.NET](#).

[What CAT Does](#)
[CAT Features](#)
[Praise for CAT](#)

What can you do in CAT?

Efficiently code raw text data sets
Annotate coding with shared memos
Manage team coding permissions via the Web
Create unlimited collaborator sub-accounts
Assign multiple coders to specific tasks
Easily measure inter-rater reliability
Adjudicate valid & invalid coder decisions
Report validity by dataset, code or coder
Export coding in RTF, CSV or XML format
Archive or share completed projects

What file types can CAT import?

Plain text
HTML
CAT XML
Merged ATLAS.ti coding

CAT Resources

[Raw Data Preparation Guide](#)
[ATLAS.ti Upload Preparation](#)
[Merging HUs in ATLAS.ti](#)

Have you tried DiscoverText?
Featuring the Facebook Graph & Twitter APIs

Ilustracja 3.1. Okno główne programu CAT z linkiem do panelu rejestracji

Źródło: opracowanie własne na podstawie programu CAT

Coding Analysis Toolkit Registration

Odebrane x

info@texifter.com
do mnie

angielski > polski Przetłumacz wiadomość

Wyłącz dla następującego języka: angielski x

(This is an automatically generated message, please do not reply to it)

Thank you for registering with the Coding Analysis Toolkit.

To complete your registration, please verify your e-mail by clicking the link below or copying and pasting the entire link into your web browser.

<http://cat.texifter.com/validateEmail.aspx?rcode=33A799C0A772BBBA85AA2A6A3ECC1A71D6E436A4D&e=qba123>

If you experience problems, please send email to info@texifter.com for help.

Odpowiedz lub przekaż dalej

Ilustracja 3.2. Potwierdzenie autoryzacji i utworzenia nowego konta

Źródło: opracowanie własne na podstawie programu CAT

Oczywiście, tak jak przystało na każdą szanującą się aplikację on-line, także w przypadku CAT natrafimy na takie podstawowe elementy, jak „przypomnienie hasła” (*forgot password?*) czy samouczek najważniejszych funkcji programu (*CAT Quickstart Guide*), a także link do instruktażowego filmu zamieszczonego na youtube (*CAT Tutorial*).

3.2. Czynności poprzedzające prace w programie – przygotowanie danych

Główną czynnością poprzedzającą rozpoczęcie właściwej pracy analitycznej w programie jest przygotowanie odpowiednio sformatowanego zestawu danych. CAT jest w tym względzie dość elastyczny, bowiem akceptowalne są trzy podstawowe formaty plików tekstowych, które można wykorzystać do pracy w programie. Są to odpowiednio: zwykły plik tekstowy (.txt), archiwum zip plików zapisanych z rozszerzeniem .txt (.zip), a także plik XML (.xml).

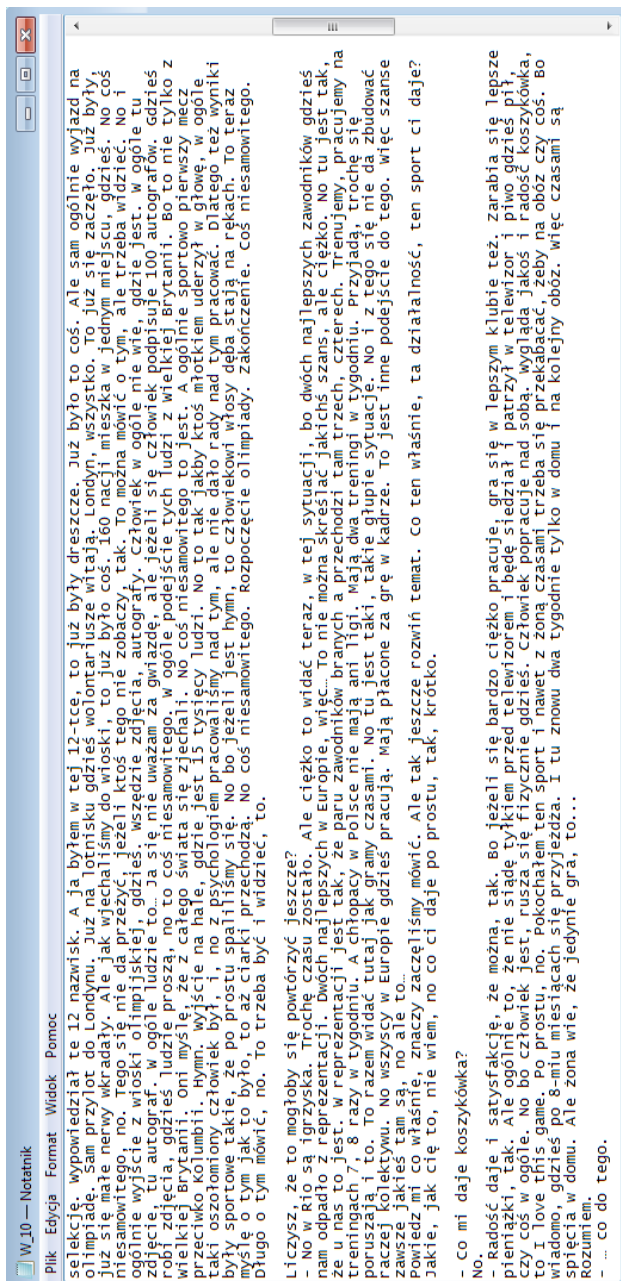
W pierwszym przypadku, a więc pliku tekstowego zapisanego w formacie .txt, program odczytuje dane jako pojedyncze jednostki podlegające kodowaniu, wówczas gdy pomiędzy tekstem zastosujemy puste wersy oddzielone od siebie za pomocą naciśnięcia klawisza enter. W związku z tym poprawne przygotowanie tego rodzaju pliku powinno polegać na zastosowaniu tak zwanych „twardych podziałów” (*hard return*), a więc rozpoczynaniu każdego nowego fragmentu tekstu, który chcemy zakodować, od nowego wiersza za pomocą każdorazowego kliknięcia klawisza enter na klawiaturze komputer. Tak przygotowany tekst będzie miał następującą postać:

```
<Tekst do zakodowania> <enter>
<enter>
<Tekst do zakodowania> <enter>
<enter> ...
```

Drugi sposób prawidłowego przygotowania danych, które mogą być wykorzystane w programie CAT, polega na użyciu archiwum zbioru plików .zip. Przy czym jeśli prześle się zarchiwizowany zbiór plików tekstowych, program odczyta każdy dokument jako jedną całość. Dlatego należy zawsze stosować jednak specjalny ogranicznik do nieprzetworzonych danych:

```
== – endcodeableunit – ==
```

Użycie takiego separatora umożliwia przesłanie archiwum zip dwóch lub więcej plików przy jednoczesnym zachowaniu możliwości kodowania różnych partii



Ilustracja 3.3. Fragment tekstu zapisanego w formacie .txt zgodnie z wymaganiami programu

Źródło: opracowanie własne na podstawie programu CAT

danych. Podobnie jak w przypadku pojedynczego pliku tekstowego, zakres tekstu do zakodowania zależy zatem od preferencji samego badacza (może to być np. zdanie, para pytań i odpowiedzi, akapit itp.).

Tak sformatowany tekst powinien mieć postać:

```
<Tekst do zakodowania> <enter>
== – endcodeableunit – == <enter>
<Tekst do zakodowania> <enter>
```

Trzeci sposób przygotowania danych polega natomiast na użyciu pliku zapisanego w formacie xml. W tym przypadku należy tak przygotować dane, aby miały one postać zgodną z poniższym schematem.

```
<?xml version = „1.0” encoding = „us-ascii” ?>
<rawcodefile>
  <codefileheader>
    <datasetname> Przykładowa nazwa zbiór danych</datasetname>
  </codefileheader>
  <allowablecodes multiplecodes = „true”>
    <code>
      <codetext>Pierwszy kod</codetext>
      <codedescription> Opis pierwszego kodu </codedescription>
      <codekey>1</codekey>
    </code>
    <code>
      <codetext>Drugi kod</codetext>
      <codekey>2</codekey>
    </code>
    <code>
      <codetext>Trzeci kod</codetext>
    </code>
  </allowablecodes>
  <items>
    <item>
      <itemtext>
<![CDATA[
```

Pierwszy fragment tekstu. Pamiętaj, aby używać znaczników CDATA dla wszystkich elementów tekstowych!

```
]]>
  </itemtext>
</item>
<item>
  <itemtext>
<![CDATA[
```

Drugi fragment tekstu. Pamiętaj, aby używać znaczników CDATA dla wszystkich elementów tekstowych!

```
]]>
  </itemtext>
</item>
<item>
  <itemtext>
<![CDATA[
```

Trzeci fragment tekstu.

```
]]>
  </itemtext>
</item>
</items>
</rawcodefile>
```

Oprócz przygotowania pliku bądź plików tekstowych (co jest niezbędne do pracy w programie) można także (opcjonalnie) utworzyć odpowiednio wcześniej pliki zawierające kody. Przy czym, jeśli ta czynność nie zostanie przez nas wykonana, wówczas nic nie stoi na przeszkodzie, abyśmy zdefiniowali kody bezpośrednio w programie, a także przypisali im dla naszej wygody odpowiednie skróty klawiaturowe.

Przygotowany wcześniej plik kodu powinien składać się z trzech części i tak jak dokumenty tekstowe musi być zapisany w formacie .txt. Treść przykładowego pliku kodu (.txt) może wyglądać następująco:

Nie | Nie, brak konkretnych nowych informacji | 1

Granica | Graniczny przypadek | 2

Tak | Tak, podano szczegółowe nowe informacje | 3

Uwaga: Treść takiego pliku musi zawierać pionowe linie |. Znajdują się one na przycisku z odwrotnego ukośnika, tuż nad klawiszem enter. Aby je uzyskać wystarczy nacisnąć shift i backslash.

Po przygotowaniu zbioru danych i pliku kodu (opcjonalnie) można przesłać je do przeglądania i do dalszego kodowania. Aby to uczynić, należy przejść do strony przesyłania nowego „surowego” zestawu danych. Znajdują się następujące informacje:

Raw Dataset Name, czyli nazwa zbioru danych, którą należy uzupełnić (jeśli tego nie uczynimy, wówczas zostanie nadana nazwa domyślna),

Raw Data File, gdzie należy po kliknięciu w przycisk *Browse* zaimportować plik z przygotowanymi danymi,

Code File, gdzie należy kliknąć na przycisk *Browse* i zaimportować plik z listą kodów,

Data Format Style, a więc określenie formatu importowanych danych. Zasadniczo są tutaj dostępne dwie opcje *Standard* (Standard) oraz *Word-by-Word* (Słowo po słowie). W większości wypadków będziemy wykorzystywać format standardowy.

Trzy pozostałe opcje, które należy wybrać na stronie przesyłania danych, to:

Pierwsza z nich to *disable verification for user-defined and multiple coding*, czyli włączanie weryfikacji zdefiniowanych użytkowników oraz możliwości wielokrotnego kodowania. Wybór tej opcji pozwala na pracę wielu użytkowników (koderów), a także na to, żeby praca zespołu przebiegała efektywnie bez powielania niepotrzebnych informacji lub dodatkowych objaśnień, wydłużających czas działania.

Allow user-defined codes, która to opcja pozwala zdefiniowanym użytkownikom na wprowadzanie nowych kodów, poza wykorzystywaniem tych, które zostały już wcześniej zaimportowane.

Allow coder to select multiple codes, zaznaczenie tej opcji umożliwia koderom wybranie więcej niż jednego kodu na dany fragment danych.

Po wybraniu odpowiednich opcji należy kliknąć na przycisk *Upload* (Załaduj). Co zakończy tę procedurę i pozwoli przystąpić do dalszej pracy.

Teraz można już edytować przesłane kody podobnie jak dodawać nowe. Po wykonaniu tego kroku należy kliknąć na przycisk *Finish* (Zakończ). W ten sposób potwierdzamy, że skończyliśmy pracę nad kodami, a to skutkuje przesłaniem zestawu danych i jego wyświetleniem w tabeli *Raw Dataset table* (tzw. surowych zbiorów danych).

View Raw Datasets for Coding

Below is the list of raw datasets you have created and uploaded:

Click on a folder name to set the project folder as the current working folder. Use the "Create New" and "Delete Folder" buttons to create and delete project folders. Note that the option to delete a project folder will only appear if the currently selected folder is at a leaf in the tree, and you cannot delete from root node. If you wish to move a dataset into a project folder, click on the dataset name in the list below and choose "Move To Folder" link in the dataset's toolbox.

Current Folder: /root node

Project Folders

[\[Create New\]](#)
—  /root node

Search for: [Search](#)

(Clear the search box and press "Search" for all rows)

Rows per page: All ▼

	Dataset	# Paragraphs	# Coders	(started/done/total)	Locked?
              	W_1	2	(0 / 0 / 0)		Unlocked
              	W_10	1	(0 / 0 / 1)		Unlocked
              	W_2	2	(0 / 0 / 0)		Unlocked
              	W_3	2	(0 / 0 / 0)		Unlocked
              	W_3-1	2	(0 / 0 / 0)		Unlocked
              	W_4	2	(0 / 0 / 0)		Unlocked
              	W_5	2	(0 / 0 / 0)		Unlocked
              	W_7	2	(0 / 0 / 0)		Unlocked
              	W_8	2	(0 / 0 / 0)		Unlocked
              	W_9	2	(0 / 0 / 0)		Unlocked

© 2007-2017 - CAT is maintained by [Texifter](#), LLC and powered by [Microsoft ASP.NET](#).

Ilustracja 3.4. Widok zestawu surowych danych (View Raw Dataset)

Źródło: opracowanie własne na podstawie programu CAT

Ostatnią czynnością, jaką powinniśmy teraz zrobić, jest przypisanie użytkowników (koderów). Należy to wykonać w kilku krokach. Przede wszystkim musimy wybrać *View Raw Datasets* (widok zestawu surowych danych), a w nim określić, jakich koderów włączymy do pracy. Przy czym jako pierwsi użytkownicy będziemy musieli utworzyć subkonta dla nowych koderów (program daje możliwość utworzenia nieograniczonej liczby subkont). Teraz klikamy na *Manage Sub-accounts*, a następnie wprowadzamy następujące informacje: nazwę użytkownika (taka, która nie jest już używana przez inne konta) oraz typ konta – tutaj możemy określić, czy będzie to regularne konto, czy też konto eksperta (są to konta przeważnie zarezerwowane dla doświadczonych badaczy, którzy mają uprawnienia wyższego poziomu w zarządzaniu danymi), informacje kontaktowe (np. nazwisko, adres e-mail itp.).

Uwaga: Każdy nowo dodany koder pojawi się na liście subkonta jako pierwszy. Osoba z subkontem otrzyma wiadomość e-mail i musi przejść proces weryfikacji przed rozpoczęciem kodowania.

Kolejnym krokiem jest przypisanie koderów do konkretnego projektu (zbioru danych). W tym celu klikamy na nazwę koderów i wybieramy opcję *Add* (Dodaj).

Uwaga: Klikając na nazwę wybranego pliku danych, zobaczymy listę posiadaczy poszczególnych kont.

Teraz klikamy na *Set Chosen Coders* w ustawieniach wybranego koderów, co spowoduje, że zostanie on dodany do listy użytkowników posiadających uprawnienia w zakresie edytowania określonej bazy danych. Dzięki temu możliwe jest śledzenie wykonywanych przez niego działań, a więc między innymi tego, ile czasu przeznaczył na pracę oraz jakich dokonał w jej trakcie zmian.

3.3. Zarządzanie uprawnieniami do zarządzania i edytowania zbiorem danych

Wszyscy użytkownicy, którzy posiadają swoje konta w ramach danego projektu (bazy danych), mogą wybierać, przeglądać i zarządzać (charakter zarządzania zależy od nadanych uprzednio uprawnień) bazą kodów dostępną w menu *Datasets*. Na stronie bazy kodów wyświetla się tabela zawierająca nazwy zestawów danych oraz m.in. informacje o tym, w jakim stopniu zostały one wykorzystane w procesie analizy.

Main Menu

Datasets

Analysis

Validation

Reports

Bookmarks

Account

Logout

Change Password

Update Account Information

Manage Sub-Accounts

Add Sub-Account

Fill in the required (denoted by a *) fields and press the "Create Account" button when complete. An e-mailed link will be sent to them with a link where they will have to click on to validate their account. Once their account is validated, they can choose their password and access to the website.

* Username:

* Required

Account Type:

Regular Account

Allow Dataset Uploads?

☐

Allow Dataset Deletions?

☐

* First Name:

* Required

Middle Initial:

* Last Name:

* Required

* Email Address:

Street Address:

City:

State/Prefecture/County:

Postal Code:

Country:

Poland

Phone:

Ilustracja 3.5. Lista użytkowników (koderów) przypisanych do wybranej bazy danych

Źródło: opracowanie własne na podstawie programu CAT

115

Upload Raw Dataset

Upload the datafile and (optional) code file below. For specific instructions on how to prepare your data and the types of datasets you can upload, [click here for a popup](#).

Raw Dataset Name *(optional)*:

Raw Data File: Nie wybrano pliku.

Suggested Code File *(optional)*: Nie wybrano pliku.

Data Format Style:

Standard ▾

☐ Disable verification for user-defined and multiple coding

☐ Allow user-defined codes

☐ Allow coder to select multiple codes

© 2007-2017 - CAT is maintained by [Texifter, LLC](#) and powered by [Microsoft ASP.NET](#).

Ilustracja 3.6. Strona Dataset Code w programie CAT

Źródło: opracowanie własne na podstawie programu CAT

Aby zalogować się jako *Primary Account Holder* (Podstawowy posiadacz konta), należy wybrać *View / Upload Raw Data* (Widok / Dodaj surowe dane) z rozwijanego menu *Dataset*. Należy kliknąć na żądany zestaw danych, co spowoduje wyświetlenie się listy dostępnych koderów, spośród których można wybrać dowolną liczbę i nadać im uprawnienia do narzędzi analitycznych w ramach projektu. Warto przy tym pamiętać, że za każdym razem, przy określaniu jakichkolwiek parametrów związanych z dostępem i uprawnieniami, trzeba je potwierdzić w celu aktywacji (*Set Permissions*).

3.4. Walidacja zakodowanych danych

Analizę niezawodności zakodowanych danych można wykonać, wybierając opcję *Standard Comparisons* (Standardowego porównania) lub *Code by Code Comparisons* (Porównania kodów). Niezależnie od tego, jaki sposób porównywania wybierzemy, należy wykonać następujące czynności. Po pierwsze trzeba określić, jakiego zestawu danych dotyczyć będzie procedura porównywania, po drugie musimy wybrać koderów, których pracę będziemy poddawać porównywaniu, i po trzecie określić, jakie konkretnie kody chcemy porównywać.

Main Menu Datasets Analysis Validation Reports Bookmarks Account Logout

Comparison Tools - Standard Compare

Dataset: -- select a dataset --

Available Coders

Add >
Add All >>
<< Remove All
< Remove

Chosen Coders

Available Codes

Add >
Add All >>
<< Remove All
< Remove

Chosen Codes

Method: Kappa [How are these calculated?](#)

☐ Include Codes that are not used in calculations
☒ Suppress Overlaps if none exist?
☐ Also show individual code pairing comparisons

Run Comparison

© 2007-2017 - CAT is maintained by [Texifter, LLC](#) and powered by [Microsoft ASP.NET](#).

Ilustracja 3.7. Analiza niezawodności kodowania

Źródło: opracowanie własne na podstawie programu CAT

W przypadku porównania standardowego zostaniemy dodatkowo poproszeni o wybranie metody porównania – Fleiss’ Kappa albo Krippendorff’s Alpha. Natomiast dla porównania kodów zostaniemy dodatkowo poproszeni, aby wybrać jeden z parametrów określających charakter porównywania: dokładne dopasowanie, nakłada się lub niedopasowanie.

Wreszcie będziemy także musieli wybrać pomiędzy jedną z opcji *Sort/Collate By* (Sortowaniem/Sortowanie przez). Po wskazaniu odpowiednich opcji należy kliknąć na *Run Comparison* (Uruchom porównanie), co spowoduje wyświetlenie się tabeli z wynikami. Wyniki można pobrać jako plik RTF.

Isaki (qbø123)



Main Menu
Datasets ▾
Analysis ▾
Validation ▾
Reports ▾
Bookmarks ▾
Account ▾
Logout

Comparison Tools - Standard Compare

Please Select a dataset

Dataset: -- select a dataset -- ▾

Available Coders

▴
▾

Add >
Add All >>
<< Remove All
< Remove

Available Codes

▴
▾

Add >
Add All >>
<< Remove All
< Remove

Chosen Coders

▴
▾

Chosen Codes

▴
▾

Method: Kappa ▾ [How are these calculated?](#)

☐ Include Codes that are not used in calculations

☒ Suppress Overlaps if none exist?

☐ Also show individual code pairing comparisons

Ilustracja 3.8. Okno z porównywarki (*Run Comparison*)

Źródło: opracowanie własne na podstawie programu CAT

Podobnie program CAT pozwala na weryfikację pracy użytkowników (koderów). Jest to proces walidacji przeprowadzonego kodowania, który polega na analizowaniu fragmentów danych i ocenie wykonanego przez poszczególnych użytkowników kodowania. Aby proces ten przeprowadzić, należy wykonać następujące czynności. Po pierwsze należy wybrać zestaw danych do sprawdzenia. Po drugie wybrać konkretny kod lub wszystkie kody istniejące w ramach zestawu danych. Po trzecie uruchomić bądź kontynuować (wcześniej rozpoczętą) procedurę sprawdzania. W każdej chwili można też zakończyć sprawdzanie, klikając linki służące do nawigacji, znajdujące się po lewej stronie.

Po uruchomieniu procedury będziemy mieli dostęp do następujących informacji: opcji filtrowania (z możliwością ich zmiany), kodu lub kodów będących podstawą sprawdzania, paragrafów, w ramach których dokonywana będzie weryfikacja, liczby pozostałych paragrafów do sprawdzenia, a także kody wybrane przez koderów oraz nazwy plików i numery akapitów.

Uwaga: Dla każdego koder będzie wyświetlała się adnotacja informująca o jego czynnościach dotyczących kodowania. Informacja o danym kodzie będzie wyświetlana pogrubioną czerwoną czcionką. Jeśli uznamy ją za poprawną, wówczas należy kliknąć na przycisk *valid* (poprawna), jeśli zaś tak nie jest, klikamy na przycisk *not valid* (niepoprawna). Ewentualnie można też nacisnąć klawisze „1” lub „Y”, jeśli adnotacja jest pozytywna lub klawisze „3” lub „N”, jeśli nie jest.

Aby zaś wyświetlić i edytować sprawdzony zestaw danych, należy: wybierać ów zestaw danych, a także koderów, kody oraz kliknąć na przycisk *edit validations*. W dalszej kolejności trzeba też sprawdzić poprawność poszczególnych kolumn, a także, czy numer podany w kolumnie kodu jest odpowiedni. Jeśli tak nie jest, należy takowy wybrać z menu rozwijanego.

Po sprawdzeniu, czy dane są kompletne, kolejną czynnością będzie wybór z górnego paska nawigacji pozycji *validations* (walidacja), a następnie *check coder statistics* (sprawdzanie statystyk dla koderów). Teraz należy jeszcze wybrać *view dataset statistics* (widok statystyk dla zestawu danych). Warto zaznaczyć, że statystyki można wyświetlić dla każdego z koderów, uzyskując w ten sposób wgląd w zakres i charakter jego pracy.

3.5. Raporty i notatki

Program CAT daje też możliwość raportowania i tworzenia notatek przydatnych w pracy analityka. Aby takie dokumenty wygenerować, należy zalogować się na swoje konto jako *Primary Account Holder* (Główny użytkownik projektu), a następnie z menu nawigacji *Reports* wybrać pozycję *Dataset Reports*. Dalej klikamy na żadaną bazę danych i dokonujemy wyboru spośród użytkowników (koderów) oraz zestawów kodów, jakie mają znaleźć się w raporcie. Ostatecznie klikamy na przycisk *Generate Report*.

Oprócz raportów dotyczących baz danych można także wygenerować raporty z wykonanej walidacji. Czynności są podobne do tych, jakie należy wykonać przy wyświetlaniu raportu o bazie danych, z tym że w tym wypadku należy z menu *Reports* wybrać pozycję *Adjudication Reports*. Reszta pozostaje bez zmian.

Wreszcie w programie CAT można także wyświetlać raporty dające wgląd w sporządzone informacje. Przy czym w tym wypadku uprawnienia takie posiadają wszyscy użytkownicy, a nie tylko ci, którzy oznaczeni są jako główni. Aby zobaczyć takie dane, należy z menu *Memo* wybrać pozycję *View Memos*, a następnie określić rodzaj filtrów tekstu lub zestawu danych oraz wskazać na sposób wysłania danych do pliku CSV, Excel lub RTF.

3.6. Zarządzanie informacjami o koncie

Przydatną opcją, w jaką wyposażony został program CAT, jest także możliwość zarządzania poszczególnymi kontami i informacjami o nich. W ramach tej opcji można dokonywać następujących zmian i modyfikacji:

Po pierwsze zmieniać hasło, co można zrobić, wybierając z menu *Account* pozycję *Change Password*. Następnie należy wpisać swoje obecne hasło, nowe hasło, a następnie potwierdzić je i kliknąć na przycisk *Change Password*.

Po drugie dokonywać aktualizacji informacji o koncie. W tym celu należy wprowadzić nowe lub zaktualizowane informacje i wybierać pozycję *Update Account*.

Po trzecie tworzyć kolejne subkonta. Aby to uczynić, trzeba wybrać opcję *Add New Subaccount* znajdującą się w górnej części strony i wypełnić wymagane pola (*), a następnie nacisnąć przycisk *Create Account* (Utwórz konto). W ten sposób zostanie wygenerowana i wysłana do nowego właściciela subkonta wiadomość e-mail z linkiem do strony, na której będzie musiał zweryfikować swoje konto. Po zweryfikowaniu konta będzie mógł wybrać hasło i uzyskać dostęp do witryny.

Po czwarte można również edytować subkonto, wybierając z menu *Account* pozycję *Manage Subaccounts*. W tym celu należy kliknąć na ikonę „Edytuj” znajdującą się po lewej stronie i w ten sposób dokonać aktualizacji stosowanych informacji.

3.7. Uwagi końcowe

Coding Analysis Toolkit (CAT) to darmowa, otwarta platforma do przetwarzania danych w chmurze, zbudowana w oparciu o technologię ASP.NET firmy Microsoft. CAT umożliwia wydajną, przejrzystą, niezawodną, prawidłową i skalowalną, opartą na sieci współpracę, dotyczącą kodowania i analizy tekstów. Jako bezpłatny, oparty na sieci Web system CAT stwarza okazję do szybkiego i wygodnego replikowania oraz łączenia danych, jak również sprawdzania wykonanych analiz przez różnych koderów. CAT ze względu na jego dostępność on-line oraz na fakt, iż jest to narzędzie bezpłatne może stanowić dogodną bazę szkoleniową w procesie nauczania.

Do niewątpliwych atutów programu należą łatwość obsługi, intuicyjność oraz przejrzysty interfejs podzielony na sześć charakterystycznych sekcji: systemu logowania i rejestracji, zarządzania kontem i subkontami, przesyłania kodów i zarządzania danymi, narzędzi porównania danych, a także narzędzi służących walidacji przeprowadzonych analiz.

O popularności i rozpowszechnieniu się opisywanego narzędzia wśród badaczy zajmujących się jakościową analizą danych najlepiej świadczą statystyki, według których istnieje obecnie 15 229 podstawowych kont CAT i 1 604 subkont. Użytkownicy CAT załadowali 9 701 zakodowanych zestawów danych i 17 065 surowych zestawów danych, a także zakodowali łącznie 2 301 451 danych, a recenzenci dokonali w sumie 209 313 weryfikacji przeprowadzonych analiz w CAT¹.

CAT jest bezpłatną usługą pierwotnie opracowaną przez Qualitative Data Analysis Program (QDAP) i był hostowany przez University Centre for Social and Urban Research na Uniwersytecie w Pittsburghu oraz QDAP-UMass w College of Social and Behaviours Sciences mieszczącym się na Uniwersytecie Massachusetts. CAT był laureatem nagrody „Best Research Software” w 2008 r. zorganizowanej przez American Political Science Association.

Zarówno sam program, jak i dodatkowe informacje o nim są dostępne na stronie: <http://cat.texifter.com/>

¹ Dane pobrane ze strony <http://cat.texifter.com/> [dostęp: 1.12.2017].

4. RQDA – narzędzie wspomagające proces analizy danych tekstowych

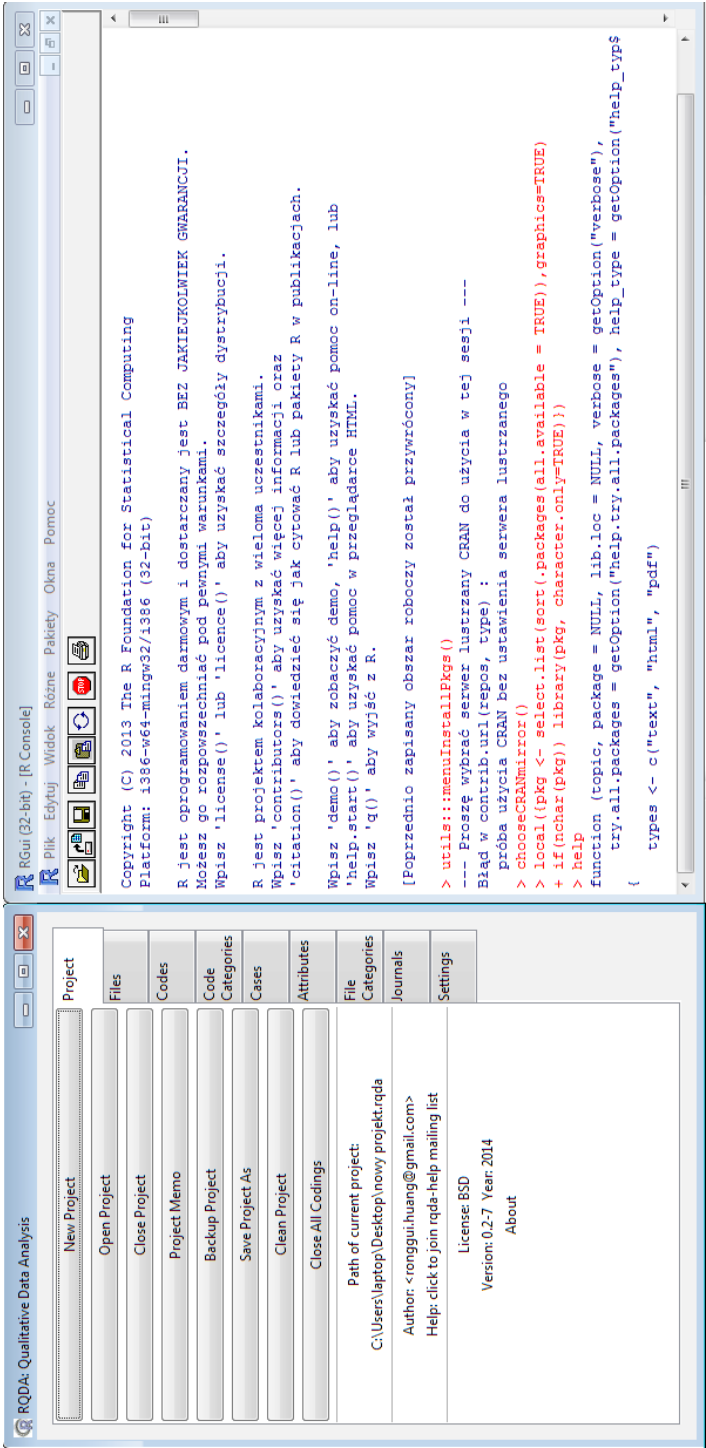
RQDA jest pakietem R do analizy danych jakościowych. Działa na platformach Windows, Linux/FreeBSD i Mac OSX. RQDA jest łatwym w obsłudze narzędziem wspomagającym analizowanie danych tekstowych. Wszystkie informacje są przechowywane w bazie danych SQLite. To program, który obejmuje wiele standardowych funkcji analizy danych jakościowych wspomaganych komputerowo. Dodatkowo integruje się z R, co oznacza, że a) statystyczna analiza kodowania jest możliwa, b) funkcje manipulacji danymi i analizy mogą być łatwo rozszerzone przez zapisywanie funkcji R. To sprawia, że w pewnym stopniu RQDA i R tworzą zintegrowaną platformę do analizy ilościowej i jakościowej danych.

4.1. Rozpoczynanie pracy w programie RQDA

Pracę w programie rozpoczynamy od utworzenia nowego projektu. W tym celu klikamy przycisk *New Project* w zakładce *Project*. W RQDA utworzony plik posiada rozszerzenie.rqda i w istocie jest bazą danych SQLite. Oznacza to, że wszystkie informacje (np. pliki, listy kodów, wszelkiego rodzaju notatki, kodowania i relacje między kodami lub między plikami itp.) są przechowywane w tym pojedynczym pliku. Zaletą tego rozwiązania jest to, że dzięki temu można bez przeszkód wykonywać kopie zapasowe oraz dokonywać migracji danych.

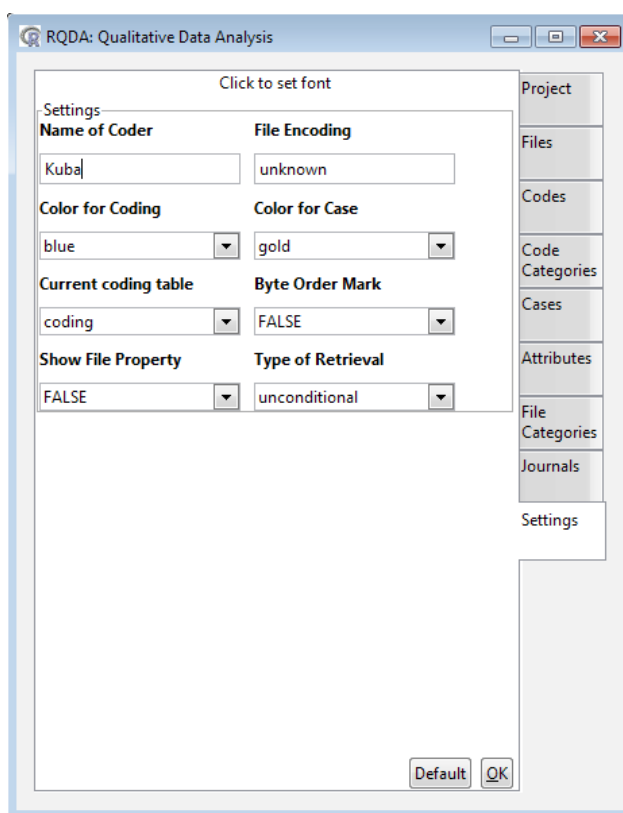
W przypadku gdy mamy już utworzony projekt i chcemy kontynuować rozpoczętą w nim pracę, wówczas powinniśmy skorzystać z zakładki *Open Project*. Dla wygody użytkownika, ale też bezpieczeństwa danych można utworzyć kopię bezpieczeństwa, wykorzystując w tym celu przycisk *Backup*. Warto przy tym zwrócić uwagę, na to, że program został wyposażony w bardzo pomocny system zapisu kolejnych wersji kopii zapasowych. Na przykład, gdy użytkownik pracuje nad plikiem demo.rqda i kliknie na przycisk *Backup*, otrzymuje plik kopii zapasowej o nazwie demo.rqda_17245808122008, gdzie kolejne cyfry oznaczają godzinę zapisu (17:24:58) oraz dzień wykonania kopii (8 grudnia 2008).

Przystępując do pracy w programie, można także dokonać personalizacji kodera (użytkownika), zmieniając domyślną wartość swoim imieniem, nazwiskiem czy nickiem. Czynność tę wykonujemy w zakładce *Settings* (Ustawienia).



Ilustracja 4.1. Okno programu z otwartą zakładką Project

Źródło: opracowanie własne na podstawie programu RQDA



Ilustracja 4.2. Widok zakładki *Settings*

Źródło: opracowanie własne na podstawie programu RQDA

4.2. Działania wykonywane na plikach (dokumentach)

Kolejną czynnością po utworzeniu bądź otwarciu gotowego już projektu jest zaimportowanie plików, na których badacz będzie wykonywał analizy. Czynność importowania wykonuje się, klikając na przycisk Importuj (*Import*) w zakładce Pliki (*Files*). Domyślnie każdy plik powinien być zapisany jako zwykły tekst z kodowaniem ASCII¹, w przeciwnym razie należy ustawić kodowanie pliku w zakładce *Settings* (Ustawienia) przed jego zaimportowaniem.

¹ ASCII (ang. *American Standard Code for Information Interchange*) – 7-bitowy kod przyporządkowujący liczby z zakresu 0–127: literom alfabetu angielskiego, cyfrom, znakom przestankowym i innym symbolom oraz poleceniom sterującym. Na przykład litera „a” jest kodowana jako liczba 97, a znak spacji jest kodowany jako 32 (Źródło: Wikipedia, <https://pl.wikipedia.org/wiki/ASCII>).

Dla wygody użytkownika wprowadzono także opcję importowania wielu plików naraz. Jeśli chcemy wykonać tę czynność, powinniśmy użyć polecenia `write.FileList`. Dzięki niemu możemy dodawać wszystkie pliki wybrane z katalogu do naszego projektu.

Zapis polecenia:

```
write.FileList(FileList, encoding = .rqda$encoding, con = .rqda$qdacon, ...)
addFilesFromDir(dir, pattern = „*.txt$”)
```

Wyjaśnienie poszczególnych Argumentów:

`FileList` – lista, gdzie każdy z jej elementów stanowi zawartość pliku, a ich nazwy `names (FileList)` są odpowiednią nazwą pliku.

`encoding` – nie należy zmieniać tego argumentu.

`con` – nie należy zmieniać tego argumentu.

`dir` – ścieżka katalogu, w którym znajdują się pliki.

`pattern` – argument, który oznacza, że importowane są tylko pliki pasujące do danego wzorca.

Więcej informacji o korzystaniu z polecenia `Write.FileList` można uzyskać m.in. na stronie: <https://www.rdocumentation.org/packages/RQDA/versions/0.3-0/topics/write.FileList>

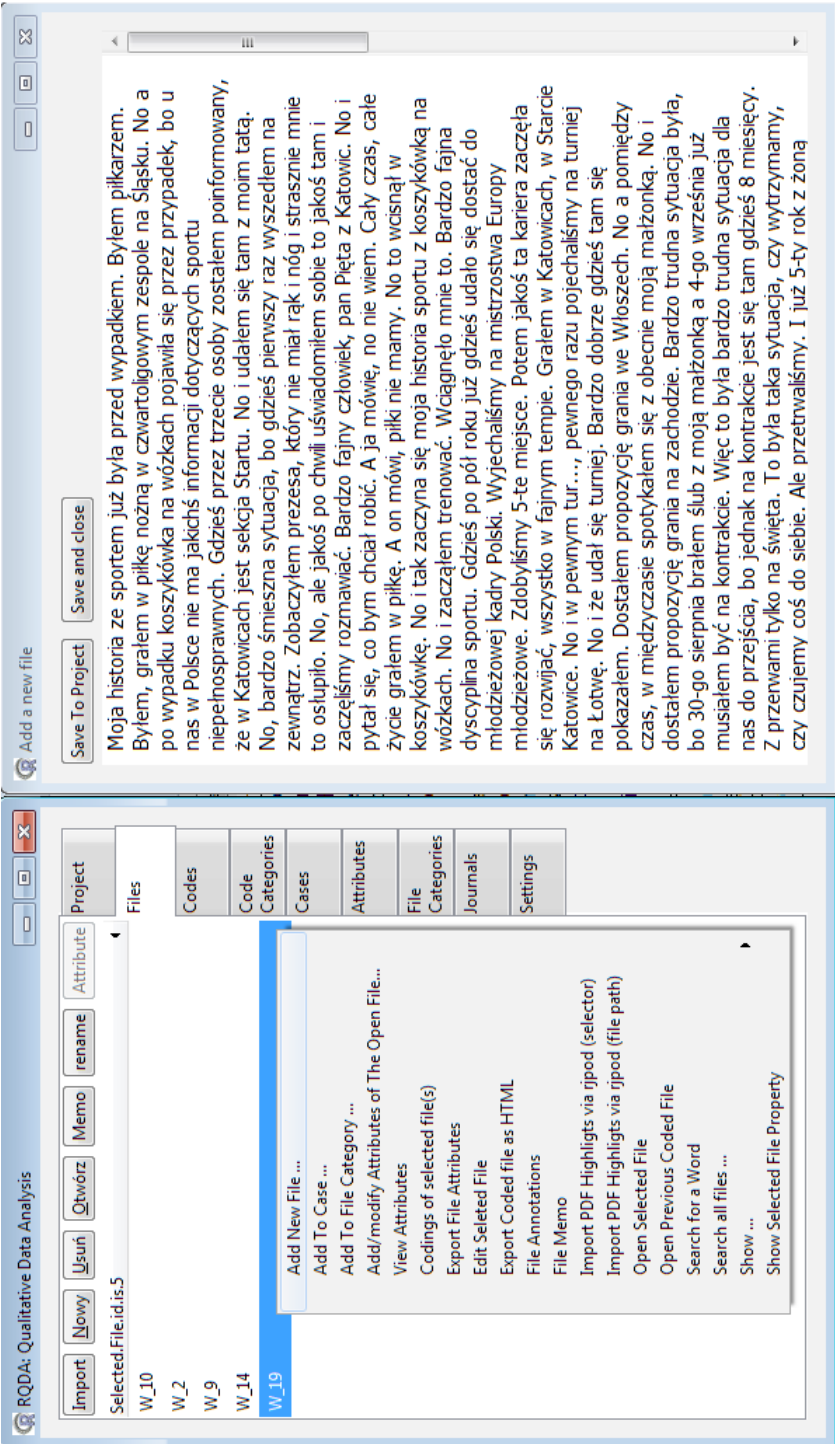
Ilustracja 4.4. Zapis oraz wyjaśnienie argumentów polecenia `write.FileList`

Źródło: opracowanie własne na podstawie programu RQDA

Ponadto korzystając z menu podręcznego, można dodać nowy plik, bezpośrednio wpisując lub wklejając jego zawartość. W tym celu należy kliknąć prawym przyciskiem myszy (gdy kursor znajduje się w polu wyświetlanej listy zaimportowanych plików), co spowoduje wyświetlenie menu kontekstowego, a następnie wybrać opcję *Add New File...* (Dodaj nowy plik...). Po wpisaniu lub wklejeniu zawartość klikamy na *Save To Project* (Zapisz w projekcie) i wpisujemy nazwę pliku oraz zatwierdzamy wykonaną operację przyciskiem OK.

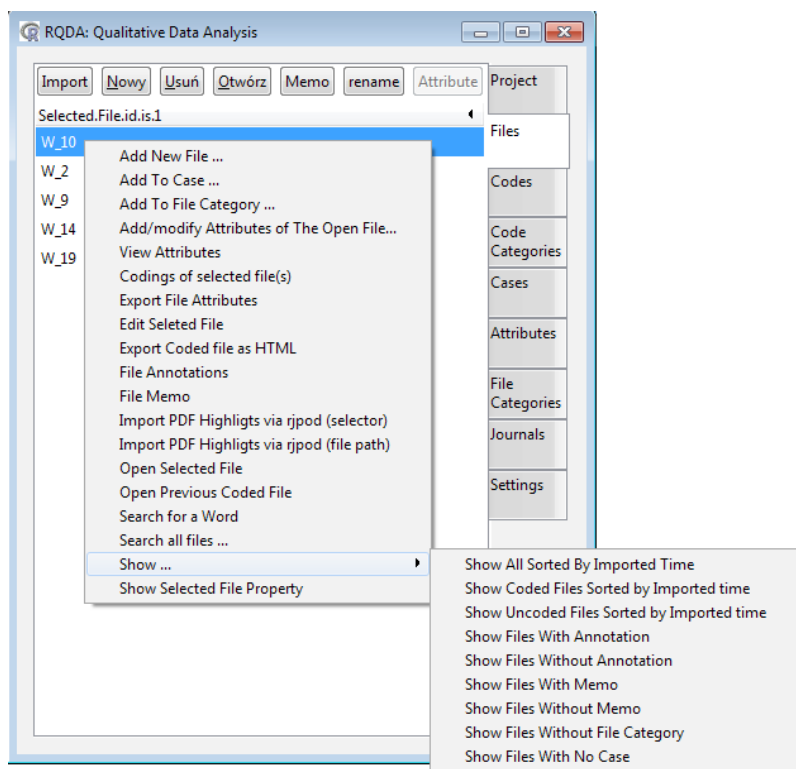
W programie RQDA możemy także swobodnie usuwać pliki oraz zmieniać ich nazwy. Aby to uczynić, należy najpierw wybrać dany plik, a następnie w zakładce *File* kliknąć na opcję *Usuń*. Z kolei, aby wyświetlić zawartość żądanego pliku, należy wybrać ów plik i kliknąć na przycisk *Otwórz*. Te sam efekt osiągniemy, klikając dwa razy na nazwę danego pliku.

Warto zaznaczyć, że w menu podręcznym można wyświetlić wszystkie rodzaje plików lub tylko dany ich podzbiór. Na przykład, można wyświetlać tylko pliki niekodowane lub wyłącznie pliki kodowane. Można także wyświetlić podzbiór plików dopasowanych do kryterium wyszukiwania według rozwijanej listy menu kontekstowego *Show...* (Pokaż).



Ilustracja 4.5. Dodawanie nowego pliku metodą wklejania i bezpośredniego wpisywania treści

Źródło: opracowanie własne na podstawie programu RQDA



Ilustracja 4.6. Menu kontekstowego *Show...* z listą możliwych sposobów wyświetlania danych

Źródło: opracowanie własne na podstawie programu RQDA

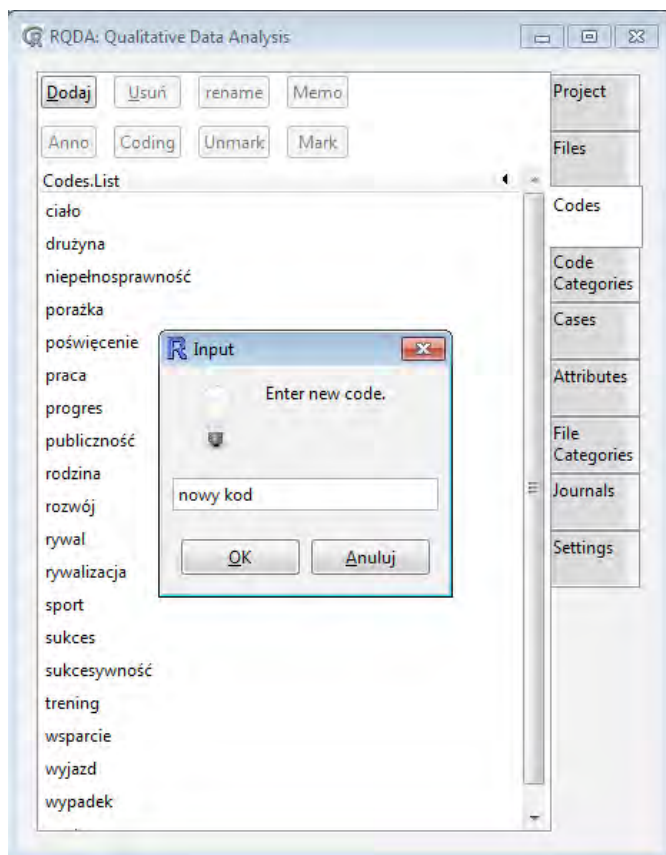
Inną funkcją, jaką oferuje program, jest sortowanie plików według ich nazw. Ponadto można obliczyć całkowitą liczbę plików w projekcie przy użyciu polecenia: `R (GetFileId)`, całkowitą liczbę niekodowanych plików przy użyciu komend (`GetFileId (type = „uncoded”)`) oraz liczbę kodowanych plików za pomocą polecenia (`GetFileId (type = „coded”)`)².

4.3. Tworzenie listy kodów i kodowanie danych

Zasadniczo istotą analizy jakościowej jest zastosowanie etykiet do segmentów tekstu. Etykieta jest zazwyczaj znaczącą koncepcją w badaniach, a więc kodem. W związku z tym rolą analityka jest wygenerowanie kodów, które można dodawać do listy, klikając przycisk *Add* znajdujący się w zakładce *Codes*. Z technicznego punktu widzenia taka czynność polega na wprowadzeniu danego kodu

² Opcje te opisano w dalszej części rozdziału poświęconego programowi RQDA.

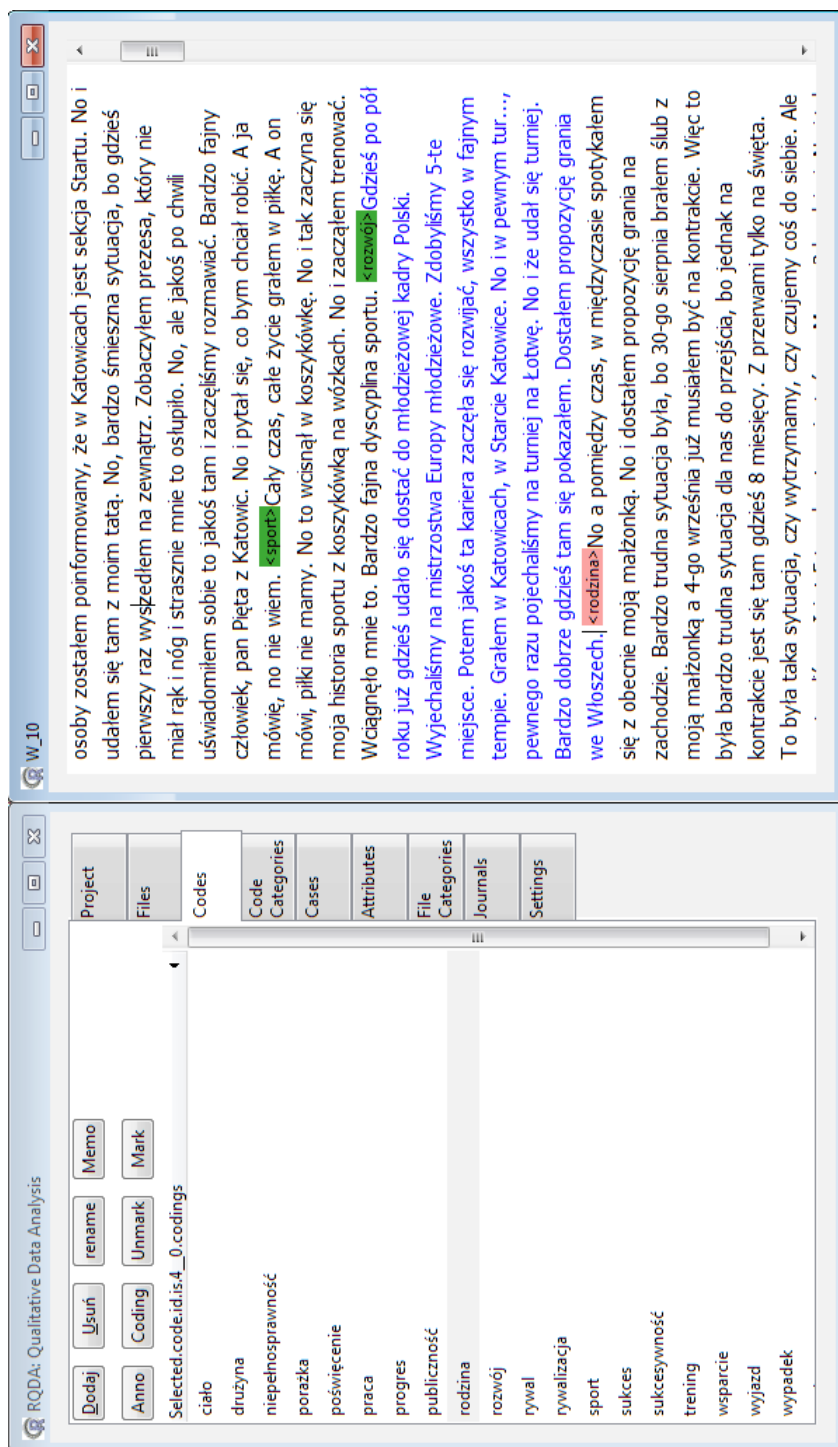
w odpowiednim oknie i kliknięciu na przycisk OK. Możesz dodać nieograniczoną liczbę kodów, ale warto zastanowić się, jakie kody są przydatne w badaniach, aby uniknąć sytuacji „przytłoczenia” ilości kategorii. Każdy kod można w dowolnym momencie edytować, zmieniając jego nazwę, a także można go usunąć z listy kodów. Czynności te są podobne do operacji na plikach.



Ilustracja 4.7. Tworzenie listy kodów w programie RQDA

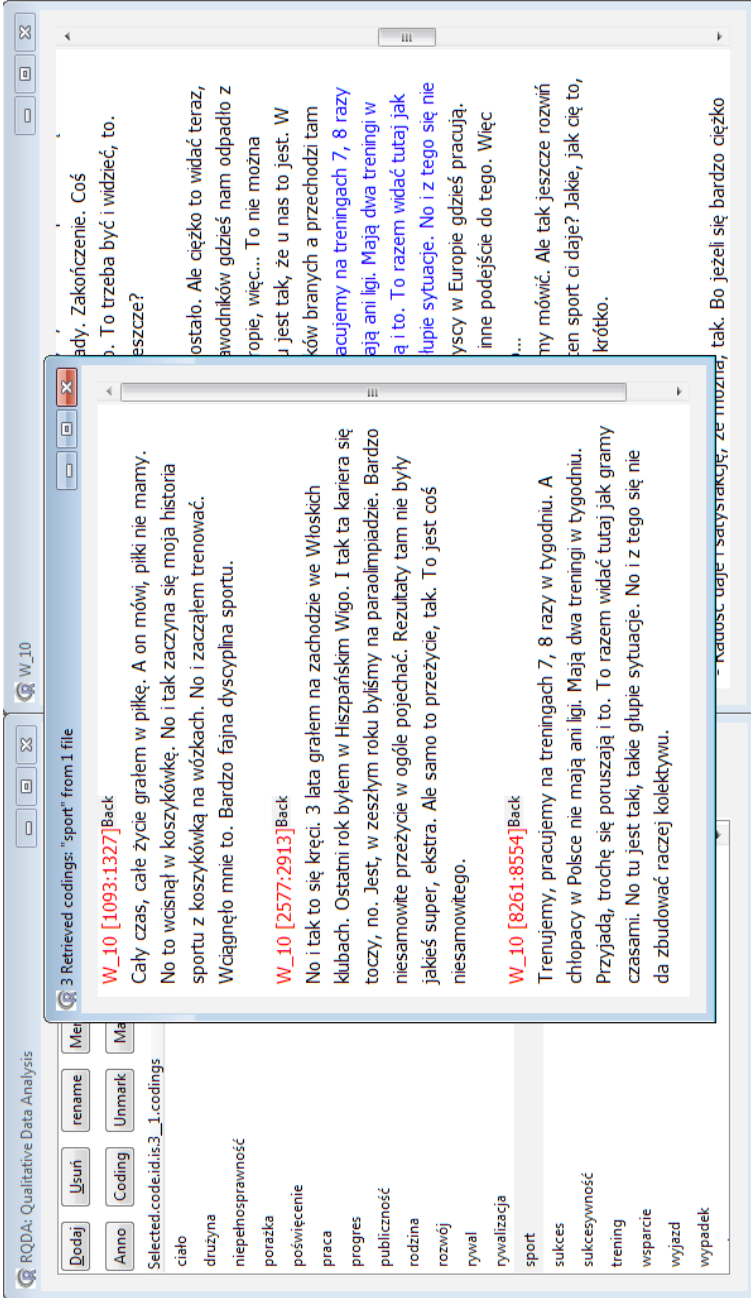
Źródło: opracowanie własne na podstawie programu RQDA

Po zaimportowaniu plików lub wygenerowaniu listy kodów można przejść do procesu kodowania. Najpierw należy otworzyć konkretny plik, a następnie przyporządkować wybrany kod do danego fragmentu tekstu. Kolejne czynności polegają w tym wypadku na: wyborze kodu, wyborze segmentu tekstu, a następnie kliknięciu przycisku znacznika (*Mark*) w zakładce Kodów. Czynność ta spowoduje, że wybrany segment tekstu zostanie wyróżniony kolorem (domyślnie jest to kolor niebieski).



Ilustracja 4.8. Kodowanie danych w programie RQDA

Źródło: opracowanie własne na podstawie programu RQDA



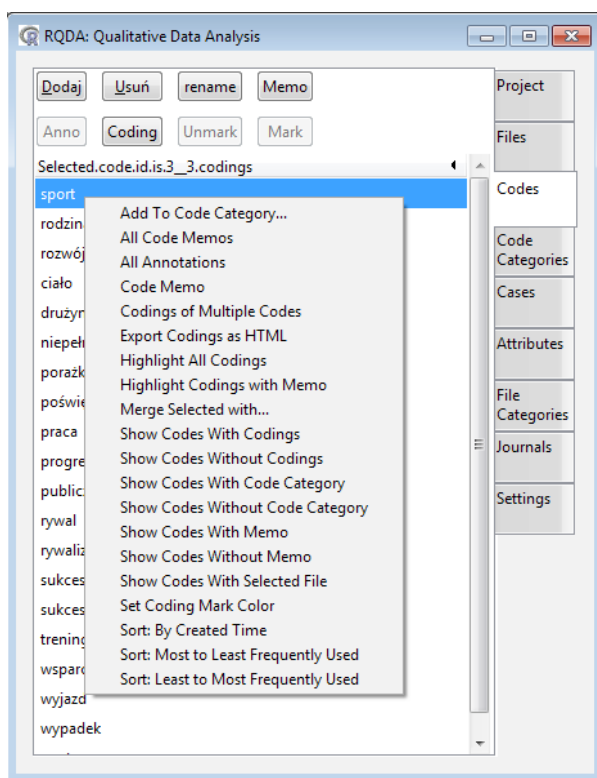
Ilustracja 4.9. Pobieranie informacji o segmentach tekstu(ów) zakodowanych wybranym kodem

Źródło: opracowanie własne na podstawie programu RQDA

Można również cofnąć kodowanie poprzez wybór tego samego segmentu tekstu i kliknięcie przycisku usuwania zaznaczenia (*Unmark*) znajdującego się w zakładce Kody.

Przydatną funkcją jest możliwość sprawdzenia, jakie informacje zostały danym kodem oznaczone. Informacje te można uzyskać, klikając dwukrotnie na dany kod.

Program pozwala także na wgląd w ogólne statystyki kodowania, co jest możliwe poprzez szereg opcji dostępnych w rozwijanym menu kontekstowym, które pojawia się po kliknięciu prawym przyciskiem myszy w dowolny kod. Dzięki temu można między innymi sprawdzić: ile razy użyto danego kodu, jaka jest średnia liczba słów zaszeregowanych do każdego kodu, czy liczbę plików powiązanych z danym kodem, a także wiele więcej przydatnych opcji.



Ilustracja 4.10. Menu kontekstowe z możliwymi operacjami na kodach

Źródło: opracowanie własne na podstawie programu RQDA

Można również pobierać zakodowane uprzednio fragmenty tekstu. Pobieranie kodów powiązanych z konkretnymi przypadkami lub zbiorami plików jest możliwe dzięki określeniu typu wyszukiwania.

4.4. Pisanie not i sporządzanie notatek w programie RQDA

Pisanie tego, co badacz myśli podczas czytania i kodowania, jest integralną częścią analizy jakościowej. Dlatego wychodząc naprzeciw tym działaniom, RQDA oferuje różnego rodzaju notatki. W przypadku pomysłów odnoszących się do całego projektu należy napisać tak zwane memorandum projektowe, wykorzystując do tego celu przycisk *project memo* znajdujący się w zakładce *Project*.

Można również napisać notatkę dotyczącą określonego pliku, klikając przycisk *Memo* znajdujący się w zakładce *Files*.

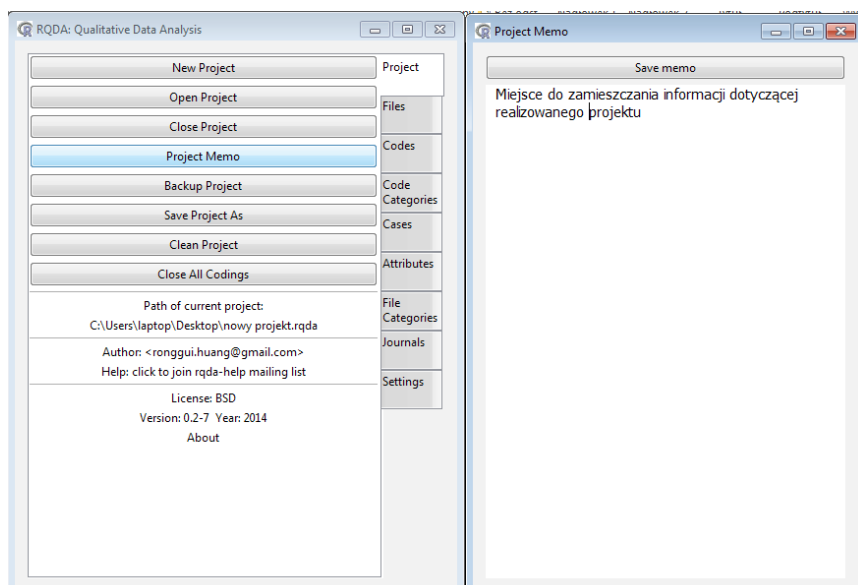
Opisami można również opatrywać poszczególne kody, stosując do tego celu przycisk *Memo* znajdujący się w zakładce *Codes*, który służy do dołączania notatki do wybranego kodu. W tym celu należy wybrać dany kod, a następnie kliknąć przycisk *Memo*, aby wyświetlić lub dodać notatkę.

4.5. Systematyzacja i zarządzanie plikami

Program został wyposażony w przydatny system zarządzania plikami. Dzięki niemu użytkownik uzyskuje dostęp do narzędzia, które pomaga organizować importowane dokumenty. Na przykład podczas analizowania raportów gazetowych możemy sklasyfikować poszczególne pliki przypisane do danego źródła (a więc określić, z jakiej pochodzą gazety). Zadanie to będzie ułatwione, jeśli utworzymy odpowiedni system „katalogów”, w ramach których będziemy gromadzić tytuły poszczególnych gazet.

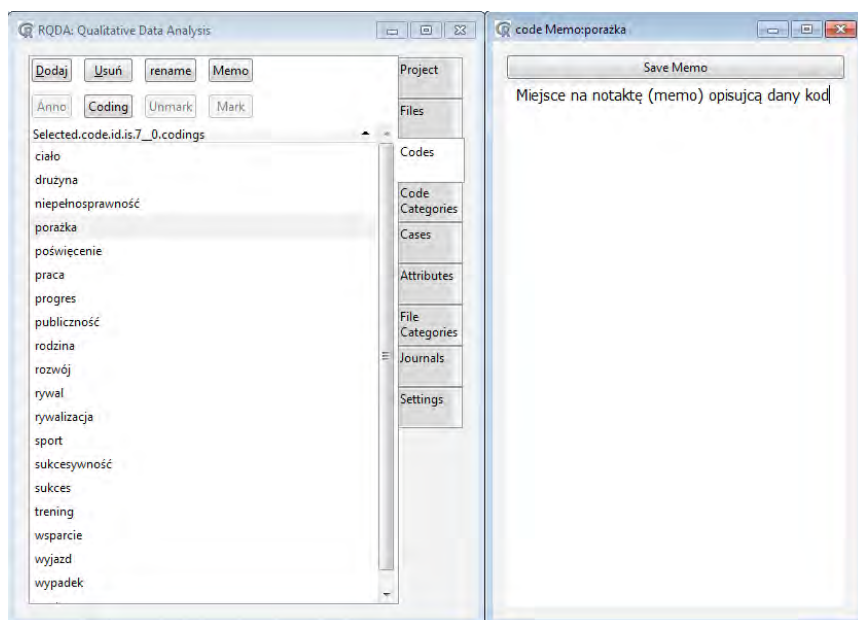
Aby to zrobić, należy kliknąć przycisk *Add* (Dodaj) znajdujący się w zakładce *File Categories*. Spowoduje to pojawienie się okna dialogowego, w którym należy wpisać nazwę tworzonego katalogu. Następnie ustawiamy kursor na nazwie nowo powstałego katalogu i klikamy raz prawym przyciskiem myszy (musi być podświetlony). Teraz można do niego dodać pliki, co robimy, klikając na przycisk *Add To* (Dodaj do). Wówczas pojawi się lista plików, z której należy wybierać właściwy i kliknąć „OK”. Równie łatwo można usunąć wybrany plik z katalogu. W tym celu należy wybrać żądany katalog, a następnie dany plik i kliknąć przycisk *Drop-From*.

Oczywiście można zmienić nazwę folderu, a także usunąć cały folder, przy czym takie działanie nie będzie miało wpływu na oryginalne pliki znajdujące się w wybranym katalogu.



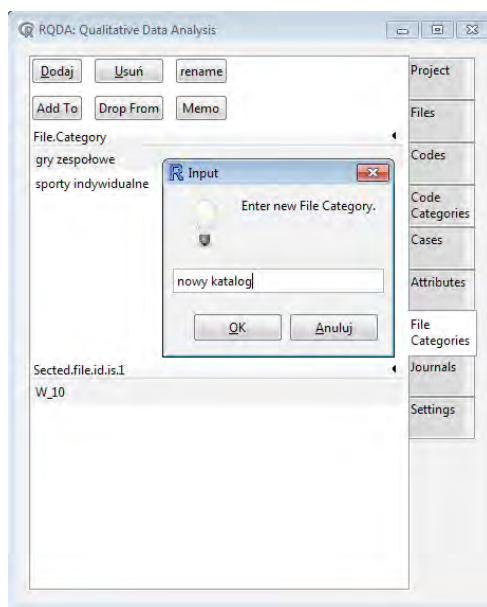
Ilustracja 4.11. Okno do umieszczania notatki (memo) w programie RQDA

Źródło: opracowanie własne na podstawie programu RQDA



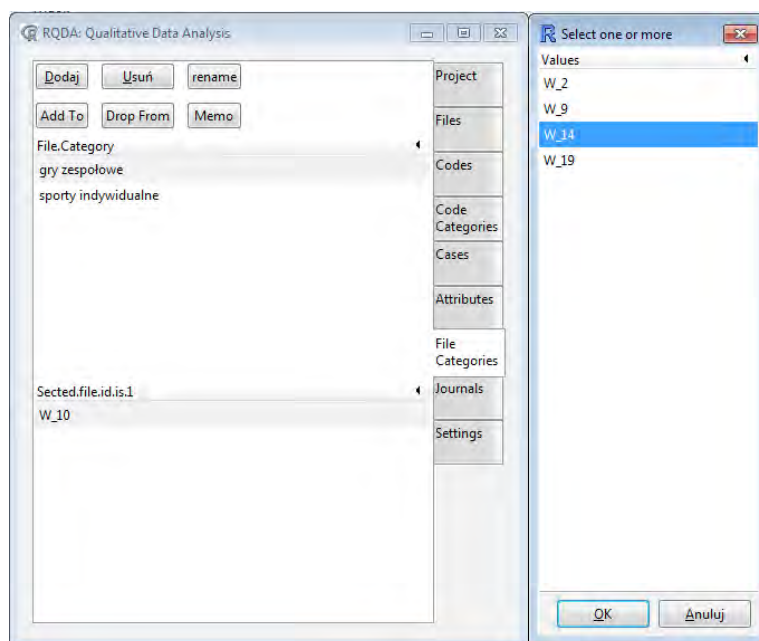
Ilustracja 4.12. Notatka opisująca wybrany kod

Źródło: opracowanie własne na podstawie programu RQDA



Ilustracja 4.13. System katalogów do zarządzania plikami

Źródło: opracowanie własne na podstawie programu RQDA

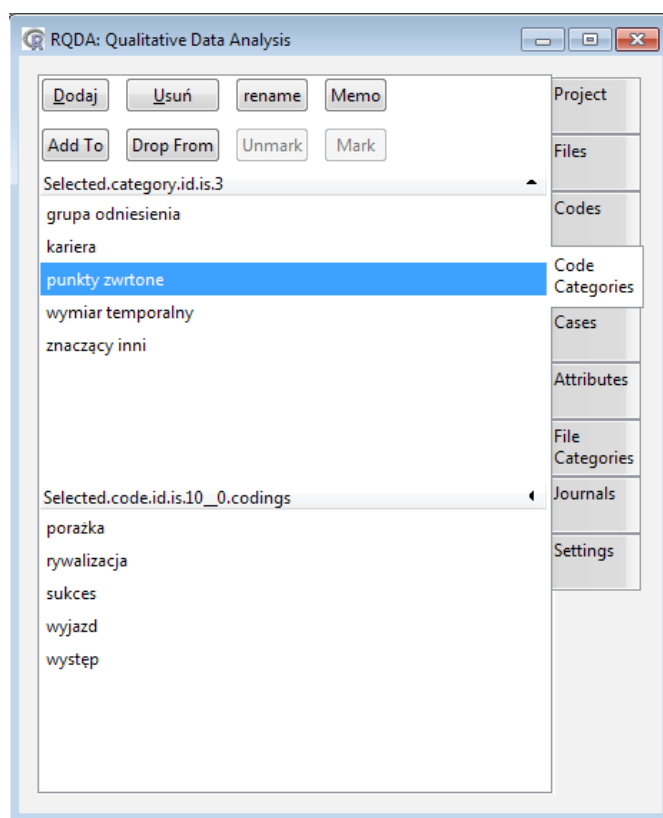


Ilustracja 4.14. Dodawanie pliku do katalogu znajdującego się w zakładce *File Categories*

Źródło: opracowanie własne na podstawie programu RQDA

4.6. Porządkowanie i zarządzanie kodami

Kategoria kodowa jest istotnym czynnikiem ułatwiającym rozwój teorii. Dwa najważniejsze kroki w tym procesie to: po pierwsze przemieszczanie się pomiędzy pojęciami o różnym poziomie abstrakcji; po drugie organizowanie tych pojęć w zależności od logicznego związku między nimi. Ten ostatni rodzaj pracy umożliwia narzędzie figurujące pod nazwą *Code Categories* obecne w programie RQDA.

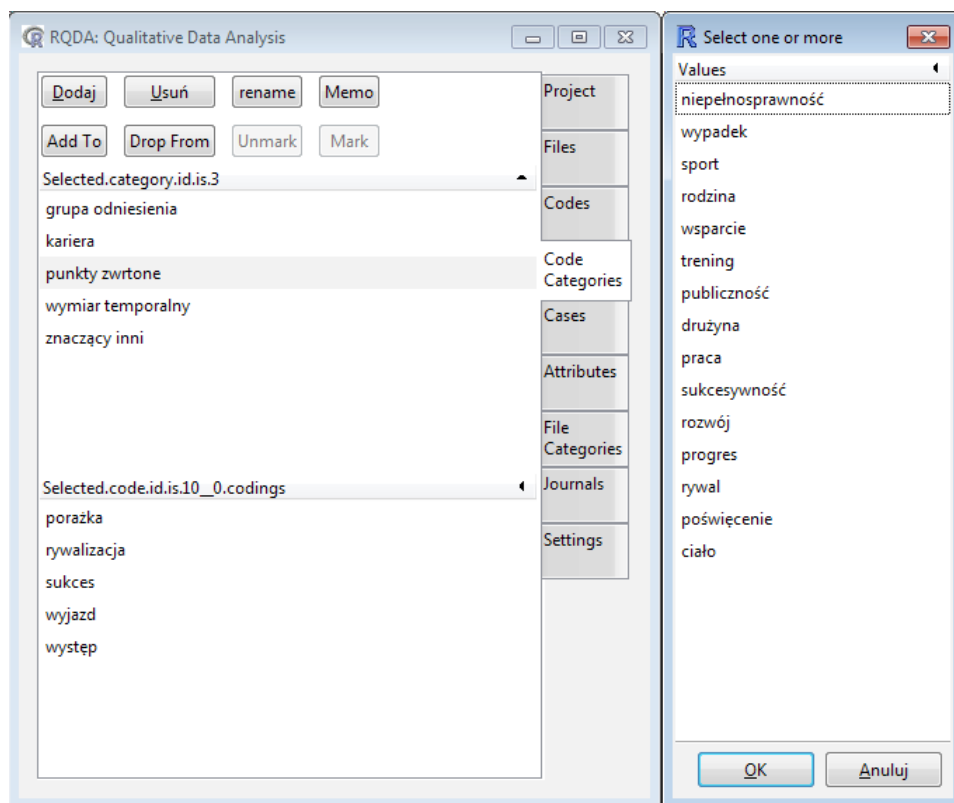


Ilustracja 4.15. Zakładka *Code Categories* w oknie programu RQDA

Źródło: opracowanie własne na podstawie programu RQDA

Używając tego narzędzia, można dodawać (przechodzić na) wyższe poziomy koncepcyjne, tworząc strukturę zbudowaną z kodów i kategorii. W tym celu wystarczy użyć przycisku *Add* (Dodaj), co spowoduje otwarcie małego okna dialogowego, w którym należy wpisać nazwę nowo tworzonej kategorii. Kolejny krok polegał będzie na dodaniu kodów do danej kategorii. Z technicznego punktu

widzenia, aby to zrobić, należy najpierw wybrać dany poziom kategorii w strukturze, a następnie kliknąć na przycisk *Add To* (Dodaj do). Spowoduje to pojawienie się listy kodów, spośród których można wybrać te, które chce się dodać, co należy ostatecznie potwierdzić, klikając na „OK”.



Ilustracja 4.16. Porządkowanie kodów w ramach tworzonej struktury kategorii

Źródło: opracowanie własne na podstawie programu RQDA

Ponieważ program RQDA charakteryzuje się dużą elastycznością, stąd naturalną sytuacją jest to, że w ramach tworzonych struktur i poziomów kategorii można dokonywać modyfikacji polegających na ich przemieszczaniu czy usuwaniu. Jeśli zatem chcemy wykonać czynność usunięcia wybranych kodów, należy najpierw wskazać na daną kategorię, co spowoduje, że zawarte w niej kody zostaną wyświetlone. Kolejnym krokiem jest wybór kodów, a następnie kliknięcie na przycisk „Drop From”.

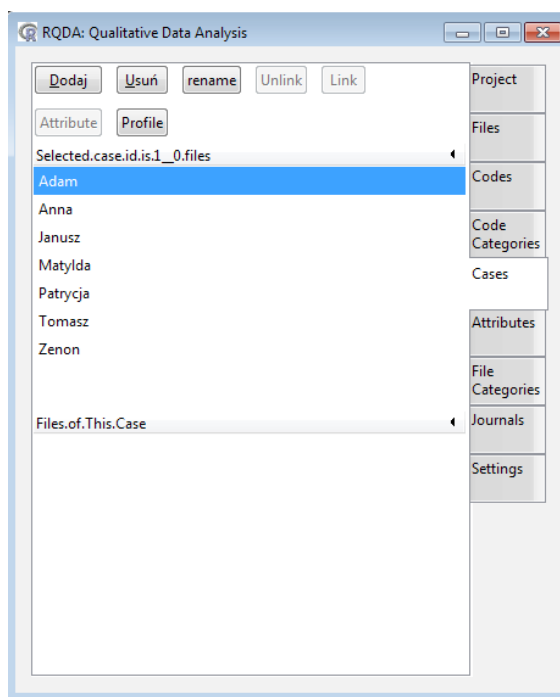
Przy tej okazji należy zauważyć, że RQDA nie dostarcza narzędzi, które pozwalałyby na utworzenie struktury przypominającej drzewo kodów i kategorii. Według stanowiska twórców programu jest to działanie celowe i planowane,

bowiem tworzenie wspomnianej struktury drzewa nie tylko nie ułatwia pracy analityka, ale wręcz ją utrudnia i niepotrzebnie gmatwa. Argumentem popierającym to stanowisko jest, że kody zazwyczaj przyporządkowuje się do więcej niż jednej kategorii, co niejako przeczy umieszczaniu ich w strukturze drzewa (z możliwością wyboru tylko jednej kategorii dla danego kodu).

4.7. Przypadki

Pojęcie przypadku (*case*) zostało wprowadzone w programie RQDA zgodnie z wymogami, jakie stawia się w związku z Jakościowym Porównywaniem Przypadków (*Qualitative Comparative Analysis*).

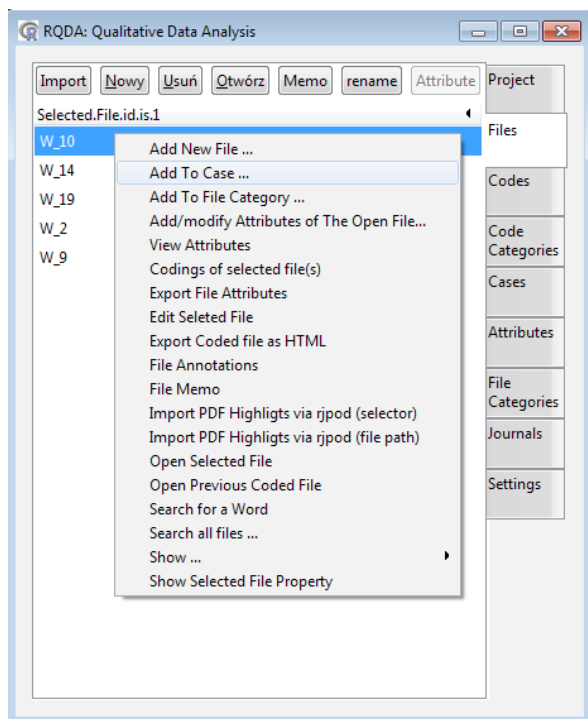
Aby przybliżyć charakterystykę przypadków, należy podać dwie zasadnicze różnice, jakie istnieją między nimi a kategoriami. Po pierwsze do przypadków można przypisać atrybuty, dzięki czemu możliwe jest wykonanie czynności ich porównywania, a po drugie można powiązać część wybranego pliku z danym przypadkiem, natomiast do danej kategorii daje się przypisać wyłącznie cały plik.



Ilustracja 4.17. Przypadki w programie RQDA

Źródło: opracowanie własne na podstawie programu RQDA

Z technicznego punktu widzenia przypisanie danego pliku do przypadku może odbywać się na dwa sposoby. W pierwszym z nich należy z zakładki *Files* wybrać żądany plik, a następnie prawym przyciskiem myszy wywołać menu kontekstowe, w którym klikamy na opcję *Add to a case...* (Dodawania przypadku).



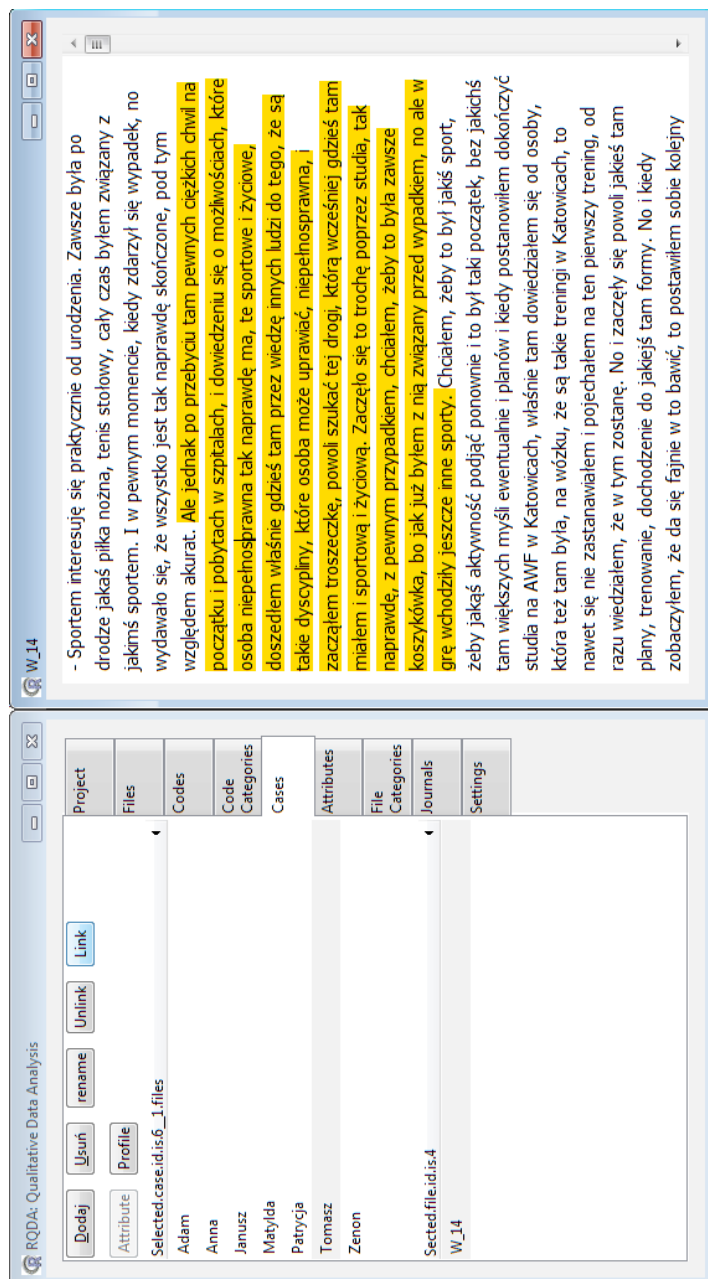
Ilustracja 4.18. Dodawanie pliku do przypadku

Źródło: opracowanie własne na podstawie programu RQDA

Użycie tej metody spowoduje dodanie całego pliku do danego przypadku. Dlatego jeśli chcemy połączyć wybraną część pliku z określonym przypadkiem, wówczas trzeba wykonać następujące kroki: otworzyć żądany plik, zaznaczyć jego fragment, wybrać dany przypadek, a następnie kliknąć na przycisk *Link*.

4.8. Atrybuty

Atrybuty w badaniach jakościowych można w zasadzie utożsamiać z tą samą koncepcją zmiennej, z jaką mamy do czynienia w badaniu ilościowym. W związku z tym, gdy „jednostka analizy” jest jasno zdefiniowana, wówczas każda występująca

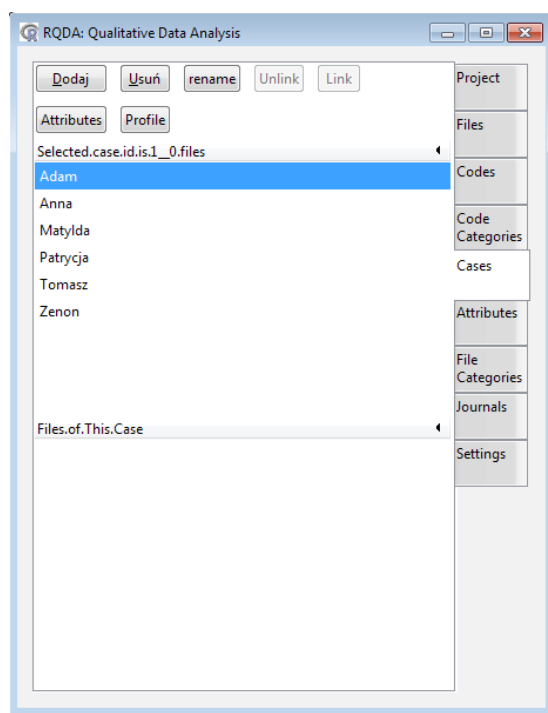


Ilustracja 4.19. Łączenie fragmentu wybranego pliku z danym przypadkiem

Źródło: opracowanie własne na podstawie programu RQDA

w danym badaniu jednostka może być traktowana jako przypadek. Przy czym atrybuty można przypisać zarówno do pliku, jak i do przypadku. Wybór zależy od tego, jaki charakter ma realizowany projekt oraz jakie cele i pytania badawcze stawiamy. Zwykle atrybuty plików są ściśle związane z samym procesem gromadzenia danych, natomiast atrybuty przypadków są bardziej analityczne, ponieważ zależą one głównie od procesu analitycznego.

Na przykład, jako odrębny przypadek może być traktowana każda osoba, która wzięła udział w badaniu. A zatem poprzez przyporządkowanie atrybutów poszczególnym osobom będzie można analizować to, jakie cechy (własności) opisujące przypadki i w jakim wymiarze mają wpływ na zachowania, działania itd. badanych aktorów społecznych w ramach interesujących nas zjawisk czy procesów.



Ilustracja 4.20. Przykład użycia atrybutów w programie RQDA

Źródło: opracowanie własne na podstawie programu RQDA

Atrybuty odgrywają dwie role. Po pierwsze, mogą być używane do określania podzbioru konkretnych zmiennych, wpływających na zachowanie i/lub przebieg badanego zjawiska, a po drugie, można na ich podstawie wykonać statystyki dotyczące tych atrybutów, co pozwala na przykład na pokazanie profilu naszych rozmówców.

4.9. Wykonywanie zapytań przy użyciu SQL

Program RQDA posiada także rozbudowane możliwości w zakresie przeszukiwania danych i wykorzystania zapytań służących między innymi do „weryfikowania” hipotez oraz konstruowanej teorii. Jednak wadą programu dość mocno utrudniającą skorzystanie z tego rodzaju narzędzi jest to, że nie występują one w formie graficznego interfejsu użytkownika. W praktyce oznacza to, że chcąc ich użyć, badacz z konieczności zmuszony zostaje do posługiwania się wpisywanymi ręcznie poleceniami.

Jedno z najbardziej popularnych i najczęściej stosowanych w praktyce przeszukiwań danych typu booleanowskiego (z zastosowaniem operatorów logicznych) można przeprowadzać w oparciu o zakodowane dane. Przy czym, aby stało się możliwe, trzeba wpisać odpowiednie polecenia w oknie programu R.

Poszczególne kody są zapisywane cyframi takimi jak 1, 2, 3 itd.

Aby zobaczyć numery różnych kodów, należy podać polecenie: `GetCoding Table()` Aby zobaczyć wybrany tekst dla kodu 1, stosuje się komendę `RQDA Query(„select seltext from coding where cid=1”)`

Przykład zastosowania operatora OR (albo) w przeszukiwaniu zakodowanych treści kodami 1 oraz 2 użyć polecenia: `: RQDAQuery(„select seltext from coding where cid=1 or cid=2”)`

Inne polecenia związane z zastosowaniem operatorów logicznych są tożsame. Przy czym, aby mieć pewność co do poprawności ich użycia, zaleca się każdorazowo zapoznać się z instrukcją programu (*Help*).

Ilustracja 4.21. Przykład zastosowania poleceń typu booleanowskiego

Źródło: opracowanie własne na podstawie programu RQDA

Aby wykonać bardziej skomplikowane wyszukiwanie typu booleanowskiego, konieczne jest już jednak programowanie.

W programie RQDA operacje wykorzystujące operatory logiczne mają postać:

`%and%` (KOD 1 i KOD 2)

`%or%` (KOD 1 **albo** KOD 2)

`and %not%` (KOD 1, **ale nie** KOD 2)

Dla przykładu, aby wykonać operację polegającą na odnalezieniu fragmentów danych zakodowanych przez KOD 1 i KOD 2, należy użyć formuły: `getCodingsByOne(1) %and% getCodingsByOne(2)`

Jeśli zaś będziemy chcieli wykonać operację polegającą na odnalezieniu fragmentów danych zakodowanych przez KOD 1, ale niezakodowanych przez KOD2, wówczas powinniśmy użyć następującej formuły: *getCodingsByOne(1) %not% getCodingsByOne(2)*

Do przeprowadzenia powyższych zapytań można również użyć następujących poleceń:

```
and(getCodingsByOne(1), getCodingsByOne(2))
or(getCodingsByOne(1), getCodingsByOne(2))
not(getCodingsByOne(1), getCodingsByOne(2))
```

Program RQDA pozwala także na bardziej złożone formuły przeszukiwania, łącząc ze sobą kilka operatorów logicznych na raz. Możemy na przykład chcieć sprawdzić, jak wygląda sytuacja z zakodowanym materiałem przez KOD 1 i KOD 2, a także zakodowanych przez KOD 3, ale nie przez KOD 4. Wówczas taka formuła przeszukiwania będzie miała postać:

```
summary((getCodingsByOne(1) %and% getCodingsByOne(1)) %or%
getCodingsByOne(3) %not% getCodingsByOne(4))
```

Należy przy tym pamiętać, że *%and%*, *%or%* and *%not%* mogą odnosić się zarówno do identyfikatorów plików, nazw plików, jak i identyfikatorów przypadków i samych przypadków.

4.10. Ustawienia i praktyczne porady

Używając program, warto pamiętać o kilku istotnych kwestiach, na które zwraca uwagę sam jego autor. Po pierwsze „nazwa koder” jest nazwiskiem bądź nickiem analityka. Parametr ten powinien być zawsze ustawiony, zwłaszcza zaś, gdy w projekcie przewidzianych jest kilka osób, które będą prowadziły analizę. Po drugie, importując pliki, należy pamiętać o ich standardzie kodowania. Przy kodowaniu ASCII należy zachować ustawienia domyślne, a więc „unknown” (nieznane). Natomiast inne kodowanie wymaga odpowiedniej konfiguracji. Na przykład, należy ustawić na UTF-8, jeśli plik jest w formacie UTF-8. Po trzecie, jeśli ustawimy typ przeszukiwania (Type of retrieval) na przypadek opcję (case), to zaznaczenie konkretnego przypadku spowoduje, że tylko ten dany przypadek będzie brany pod uwagę podczas przeszukiwania. Jeśli jednak nie zostaną wybrane żadne przypadki, wówczas wszystkie istniejące przypadki będą brane pod uwagę podczas przeszukiwania danych. Tożsama sytuacja występuje podczas przeszukiwania

plików. Po czwarte w systemie Linux należy ustawić UTF-8, aby lepiej obsługiwać inne niż zapisane w języku angielskim dokumenty. Po piąte po uruchomieniu RQDA lepiej zminimalizować okno główne R, co ułatwia pracę w programie RQDA. Po szóste oprócz zastosowania funkcji służących porządkowaniu plików można użyć w tym celu listy kodów i przypadków, wykorzystując określone reguły nazywania plików. Na przykład, jeśli nazwa grupy kodów zawiera 1 jako przedrostek, a inna grupa kodów zawierająca 2 jako przedrostek, można kliknąć symbol trójkąta znajdujący się w prawym górnym rogu listy kodów (Codes List), co spowoduje, że kody zostaną posortowane według ich prefiksów. Po siódme, jeśli dany plik jest duży, a my chcemy zamknąć program bez kończenia jego kodowania, należy napisać notatkę (np. o nazwie „Niedokończone”). Następnie można użyć funkcji *Search File...* (Szukaj pliku...) dostępnej w menu podręcznym, aby zlokalizować wspomniany plik, wpisując jego nazwę „% niedokończone”.

Generalnie, dokonywanie wszelkich zmian w ustawieniach programu jest dość intuicyjne i nie powinno nastęrczać większych problemów. Jeśli zaś nie będziemy pewni co do wprowadzonych przez siebie modyfikacji, zawsze możemy wybrać opcję ustawień domyślnych („default”).

4.11. Uwagi końcowe

RQDA jest pakietem R służącym do analizy danych jakościowych. Działa on na platformach Windows, Linux/FreeBSD i Mac OSX. RQDA jest łatwym w obsłudze narzędziem wspomagającym analizowanie danych tekstowych. Wszystkie informacje są przechowywane w bazie danych SQLite za pośrednictwem pakietu RS z RSQLite. Wśród wielu opcji, które wspomagają proces analizy danych, można wyróżnić: kodowanie danych; segregowanie i porządkownie wszystkich zaimportowanych oraz utworzonych danych; przeszukiwanie zakodowanych informacji, w tym z zastosowaniem operatorów logicznych; wyszukiwanie w plikach słów kluczowych; tworzenie przypadków oraz atrybutów, co przydaje się między innymi w badaniach opartych na metodach jakościowo-ilościowych; a także wszechstronną i elastyczną możliwość edytowania i organizowania wszystkich plików.

Poza wymienionymi i niewątpliwymi atutami programu, przyglądając się krytycznie jego działaniu, można stwierdzić, że pewną wadą programu jest to, że obecnie obsługuje tylko sformatowane dane w formacie zwykłym. Ponadto nie wszystkie opcje wspomagające analizę danych są dostępne na zasadzie przycisków i okienek. W związku z tym dla niektórych z nich niezbędne jest skorzystanie z wiersza poleceń zapisywanego w programie R. Z drugiej strony dzięki funkcjom R można między innymi: importować grupę plików jednocześnie; obliczać relację między dwoma kodowaniami, biorąc pod uwagę indeksy kodowania; eksporto-

wać plików; wykonywać w sposób elastyczny i dopasowany do wymagań konkretnego badacza przeszukiwanie danych, w tym także z zastosowaniem złożonych formuł obejmujących operatory logiczne.

Ważną kwestią jest także to, że program RQDA jest nadal intensywnie rozwijany, co z pewnością przyczyni się do jego modernizacji polegającej nie tylko na wprowadzeniu nowych funkcji, ale także stworzeniu jeszcze bardziej przyjaznego środowiska dla użytkownika. Przede wszystkim zaś warto podkreślić, że program cechuje duża niezawodność i stabilność, a także to, że wszystkie dane zapisywane są w jednym pliku (projekcie) i, co warto podkreślić, dzieje się to w sposób automatyczny, bez konieczności ręcznego zapisywania danych.

Program RQDA jest narzędziem służącym do analizy danych jakościowych powstałym na bazie oprogramowania R.

Pliki instalacyjne są dostępne na stronie:

<http://rqda.r-forge.r-project.org/>

Systemy, pod jakimi można zainstalować oraz uruchomić program: Windows, Linux, Mac OS.

Proces instalacji jest zróżnicowany w zależności od rodzaju systemu operacyjnego użytkownika oraz tego, czy instalacja odbywa się po wcześniejszym zainstalowaniu programu R, czy też ma być przeprowadzona niezależnie od niej. Dokładne informacje pokazujące proces instalacji krok po kroku można znaleźć na stronie internetowej projektu RQDA.

Przy opracowywaniu materiału o programie posłużono się instrukcją użytkownika opracowaną przez HUANG Ronggui (2016). RQDA: Analiza danych jakościowych R. R w wersji pakietowej 0.2-8.

5. STORIES MATTER – program do porządkowania i archiwizowania materiałów audiowizualnych

Stories Matter to bezpłatne narzędzie dostępne na zasadzie open source, stworzone z myślą archiwizacji cyfrowych materiałów wideo i audio. Zgodnie z zamysłem jego twórców, ma ono stanowić alternatywę dla transkrypcji, a więc umożliwiać gromadzenie, ale także sortowanie i porządkowanie zgromadzonych materiałów. Co więcej, program pozwala na edytowanie nagrań według własnych kryteriów użytkownika, a także na tworzenie playlist zawierających nagrania uporządkowane według określonych tematów czy problemów badawczych. Ponadto twórcy oprogramowania wzbogacili je o przydatne funkcje przeszukiwania danych, pozwalające w szybki sposób na zlokalizowanie wybranych nagrań, sesji czy wcześniej utworzonych tematów. Użytkownicy mogą również eksportować wyniki swoich prac w kilku różnych formatach, co ułatwia ich wykorzystanie w prezentacjach czy przy projektowaniu witryn internetowych. Wreszcie warto podkreślić, że program może posiadać wbudowane funkcje on-line, które umożliwiają współpracę wielu użytkownikom przy tworzeniu jednej bazy danych.

5.1. Wygląd okna programu

Okno programu prezentuje się bardzo przyjaźnie dla użytkownika. Nie ma tu zbędnych dodatków czy niepotrzebnych ikon. Wszystko ułożone jest w sposób bardzo przejrzysty, co powoduje, że osoba korzystająca z programu niemalże od razu oswaja się z jego wyglądem.

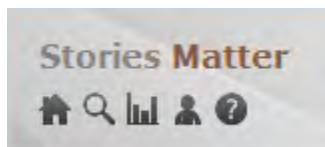
Po uruchomieniu Stories Matter naszym oczom ukazuje się okno podzielone na kilka obszarów. W lewym górnym rogu mamy więc obszar z dwiema zakładkami *Projects* oraz *Playlists*. W pierwszej uzyskujemy dostęp do poszczególnych baz danych, w których umieszczane są wszystkie elementy danego projektu. Dzięki niej możliwe jest przeglądanie całej zawartości bazy danych w jednym oknie. Przechodząc do drugiej zakładki, uzyskujemy możliwość przeglądania i otwierania istniejących zestawień zaimportowanych klipów (a więc dokumentów audiowizualnych). Zapewnia ona szybki sposób poruszania się między projektami, wywiadami, sesjami i klipami. Poniżej znajduje się obszar, w którym wyświetlane są poszczególne słowa (ujęte w tak zwanej „chmurze tagów”), natomiast prawa



Ilustracja 5.1. Okno programu Stories Matter

Źródło: opracowanie własne na podstawie programu Stories Matter

kolumna służy do wyświetlania treści odnoszących się do konkretnego pliku audio-video. Jednak najwięcej miejsca zajmuje obszar, w którym wyświetlane są bieżące materiały, do jakich odwołuje się użytkownik programu. Charakterystyczne są także ikonki znajdujące się w prawym górnym rogu okna programu, które pozwalają na dostęp do kilku najistotniejszych funkcji programu (strona główna, wyszukiwanie, eksportowanie, preferencje użytkownika i ikony pomocy).



Ilustracja 5.2. Ikonki szybkiego dostępu do wybranych funkcji programu

Źródło: opracowanie własne na podstawie programu Stories Matter

Pierwsza z ikon *Home* (Strona główna) pozwala na powrót do listy projektów zapisanych w programie. Drugą dostępną opcją jest ikona wyszukiwania *Search* (Wyszukiwanie), która umożliwia użytkownikowi przeszukiwanie w ramach istniejących projektów, wywiadów, sesji czy klipów. Kliknięcie na trzecią ikonę *Export data* (Eksportowanie) powoduje natomiast przejście do opcji eksportu, która umożliwia przeniesienie zawartości istniejącej bazy danych do postaci prezentacji PowerPointa lub dokumentów html. Z kolei czwarta ikona *User settings* (Preferencje użytkownika) odpowiada za dostęp do funkcji związanej z wyborem języka. Przy czym w przypadku programu Stories Matter będzie to wyłącznie wybór pomiędzy językiem angielskim lub francuskim. Wreszcie piąta ikona *About Stories Matter* (Pomocy) prowadzi do menu pomocy i ma dostarczyć użytkownikom podstawowych informacji na temat korzystania z oprogramowania.

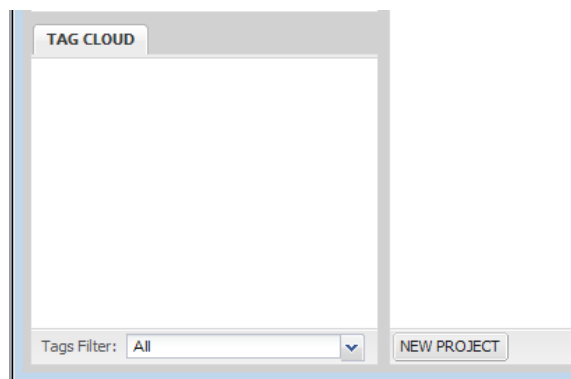
5.2. Rozpoczynanie pracy w programie Stories Matter

Na wstępie warto zaznaczyć, że architektura programu Stories Matter charakteryzuje się swego rodzaju „warstwowością” (używając słów jej twórców). Oznacza to, że w skład projektu, a więc utworzonej bazy danych, mogą kolejno wchodzić rozmówcy, sesje, klipy i playlisty.

Aby rozpocząć tworzenie wspomnianej bazy danych, należy kliknąć na przycisk *New Project* (Nowy Projekt) znajdujący się na dole okna programu.

Po wyborze wspomnianej opcji pojawi się okienko programu *Project Editor*, w którym należy wypełnić kilka pól zawierających informacje potrzebne do utworzenia projektu. Pierwsze pole wymaga, aby użytkownik nadał projektowi

tytuł (*Project Name*). W drugim polu należy wpisać krótki opis swojego projektu (*Short Description*), który później pojawi się na stronie startowej Stories Matter wśród listy projektów. W trzecim polu użytkownik musi wprowadzić dokładniejszy opis projektu (*Description*), który będzie wyświetlany po prawej stronie listy projektów. Po wypełnieniu powyższych pól użytkownik powinien wybrać i zaimportować odpowiedni obraz dla projektu (*Project Image*). Na koniec użytkownik powinien zdecydować, czy nowo utworzony projekt będzie widoczny dla współpracowników, którzy mogą mieć dostęp do serwera bazy danych. Jeśli chcemy, aby tak się stało, wówczas powinniśmy wybrać opcję *This project is public*. W ten sposób inne osoby będą mogły wyświetlać i edytować projekt. Jeśli pozostawimy wspomniane pole puste, wówczas projekt zostanie ukryty przed innymi, a użytkownik będzie mógł pracować w trybie on-line.



Ilustracja 5.3. Przycisk umożliwiający utworzenie nowego projektu

Źródło: opracowanie własne na podstawie programu Stories Matter

Po wykonaniu tych kroków należy kliknąć na przycisk *Save Project* (zapisywanie projektu). W tym momencie okno edytora projektów zostanie zamknięte, a użytkownik zobaczy, że nowy projekt został dodany do listy projektów.

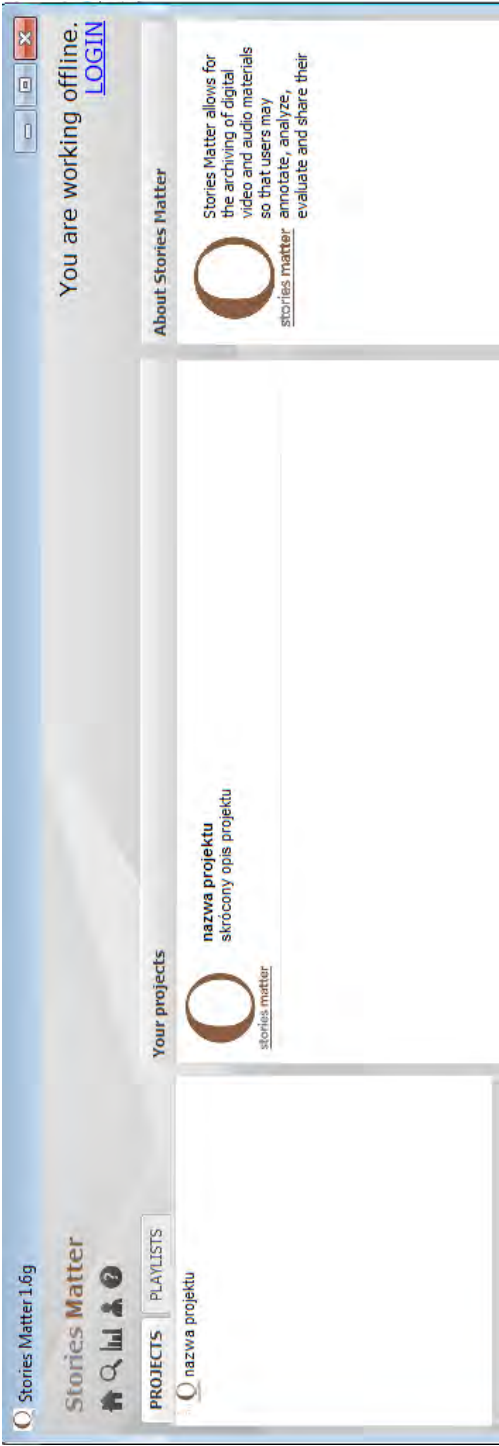
W przypadkach gdy na liście projektów znajduje się kilka pozycji, użytkownicy mogą, klikając bądź na obraz projektu, bądź na krótkim jego opisie pokazanym na liście projektów, szybko przeglądać zawartość każdego z nich. Funkcja ta pozwala użytkownikom wyświetlić szczegółowy opis różnych projektów w przestrzeni po prawej stronie okna programu.

Aby edytować lub usunąć istniejący projekt, należy wybrać stosowną opcję spośród znajdujących się na dole okna listy projektów. Przy czym trzeba pamiętać, że wybranie opcji usunięcia projektu spowoduje usunięcie wszystkich rozmówców, sesji i klipów zawartych w projekcie, a także wszystkich powiązanych z nim notatek.



Ilustracja 5.4. Okno Project Editor służące do tworzenie nowego projektu

Źródło: opracowanie własne na podstawie programu Stories Matter



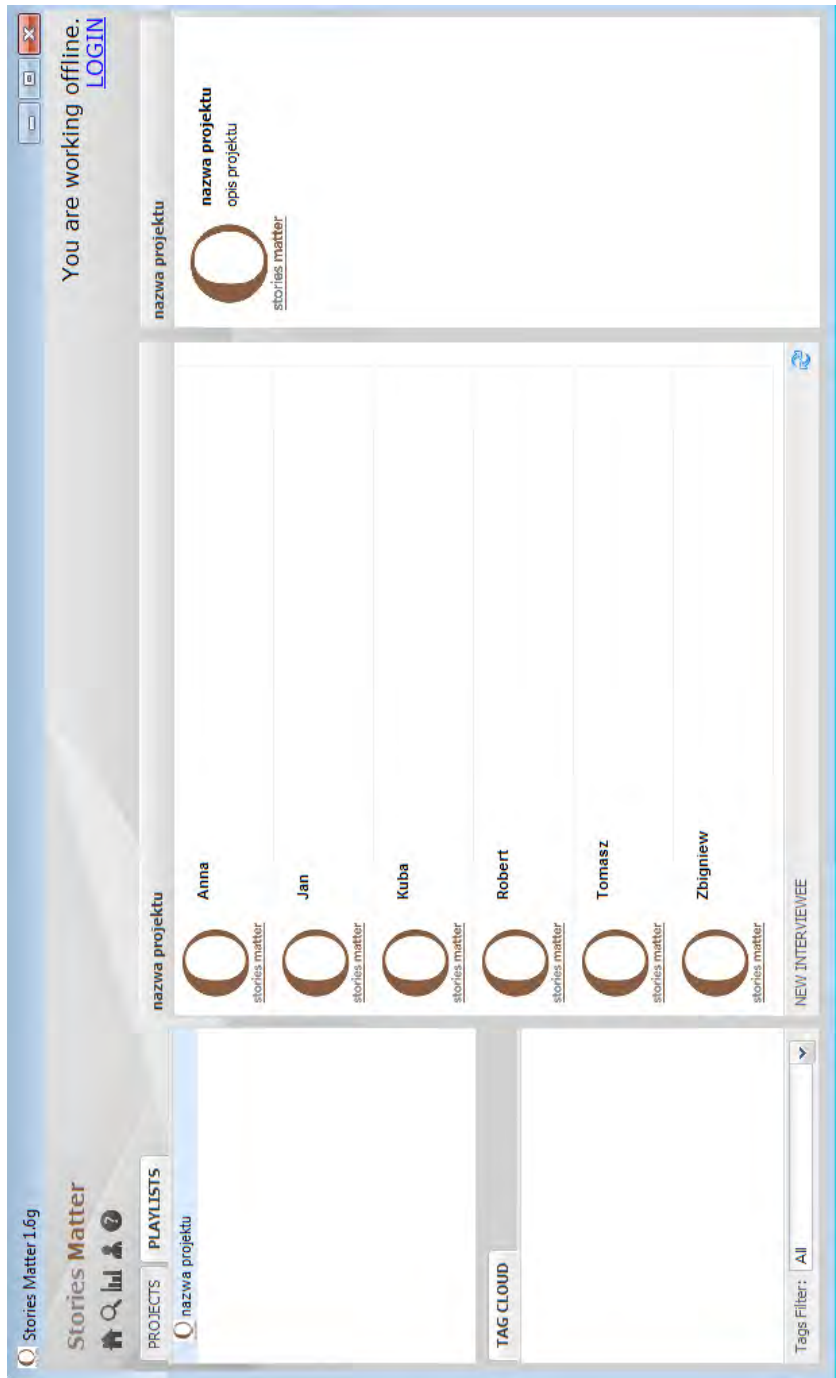
Ilustracja 5.5. Okno listy dostępnych projektów (*Project List*)
Źródło: opracowanie własne na podstawie programu Stories Matter

5.3. Tworzenie i edytowanie informacji o rozmówcach

Kolejnym poziomem czy też wspomnianą wcześniej „warstwą” jest zarządzanie i edytowanie informacji o rozmówcach. Program Stories Matter pozwala na dodawanie, modyfikowanie, usuwanie, ale też opisywanie czy oznaczanie rozmówców według potrzeb użytkownika. Oznacza to, że dzięki tej opcji zyskuje się narzędzie spersonalizowane, które umożliwia dość swobodne, a zarazem różnorodne operacje związane z przetwarzaniem informacji o rozmówcach. Lista wszystkich rozmówców zostanie wyświetlona, gdy użytkownik dwukrotnie kliknie na obrazek towarzyszący danemu projektowi lub na jego krótki opis.

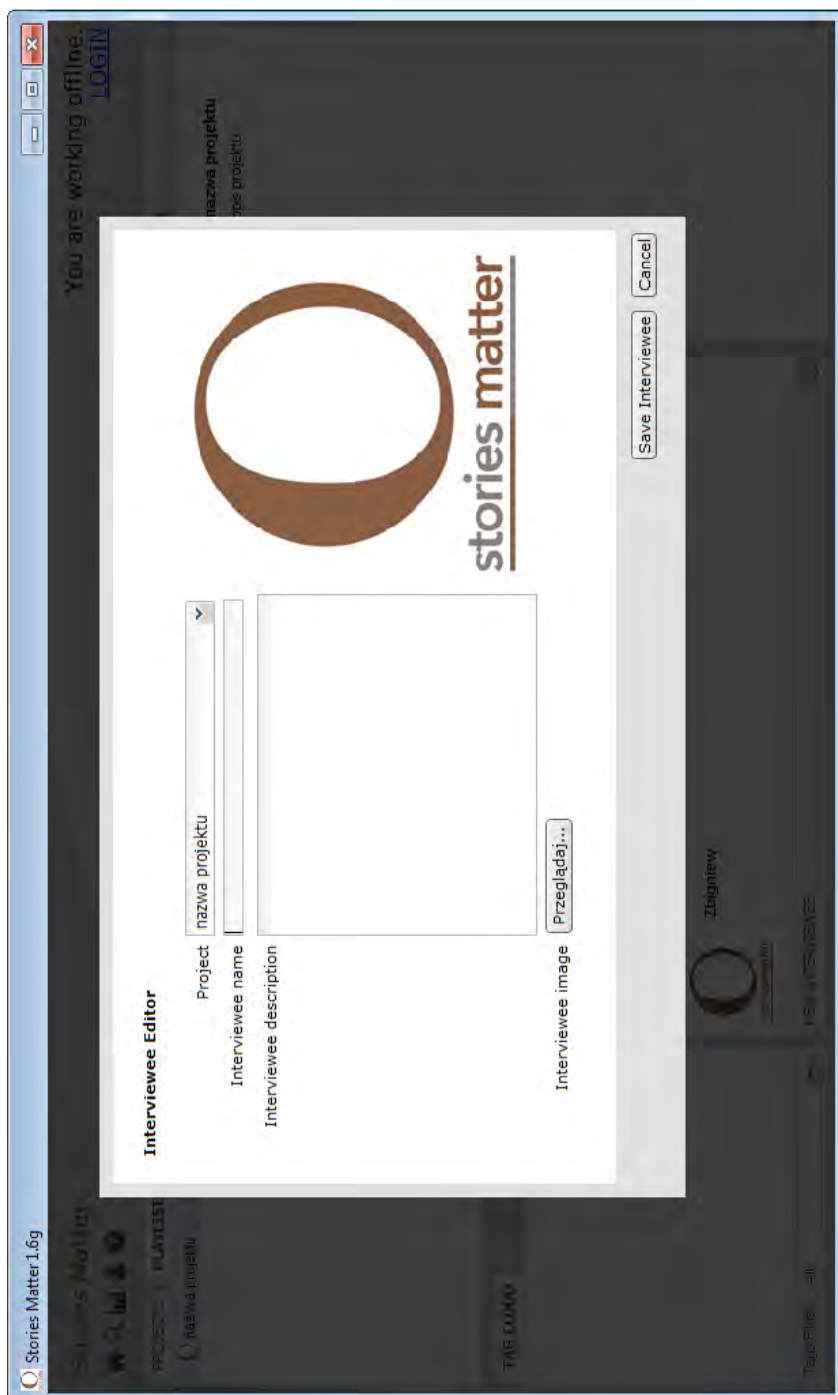
5.3.1. Podstawowe czynności związane z dodawaniem, edytowaniem i usuwaniem rozmówców

Aby dodać nowego rozmówcę, należy wybrać konkretny wywiad spośród dostępnych na liście wywiadów. Operacja ta spowoduje wyświetlenie się okienka dialogowego (*Interviewee editor*), w którym należy kolejno podać nazwę (*Interviewee name*) oraz opis (*Interviewee description*) dotyczący rozmówcy. W oknie można również dodać zdjęcie rozmówcy, przesyłając je np. ze swojego komputera. Po tych czynnościach należy jeszcze potwierdzić chęć zapisu tak wprowadzonych informacji (klikając przycisk *Save*).



Ilustracja 5.6. Lista rozmówców

Źródło: opracowanie własne na podstawie programu Stories Matter



Ilustracja 5.7. Okno edycji rozmówców

Źródło: opracowanie własne na podstawie programu Stories Matter

Klikając na nazwę lub zdjęcie rozmówcy spośród widocznych na liście, wyświetlony zostanie jego szczegółowy opis (pojawi się w oknie programu po jego prawej stronie). Jednocześnie wyświetlą się opcje służące do edycji lub usuwania poszczególnych rozmówców.

NEW INTERVIEWEE EDIT INTERVIEWEE DELETE INTERVIEWEE DUPLICATE INTERVIEWEE

Ilustracja 5.8. Pasek opcji służący do edycji informacji o rozmówcy

Źródło: opracowanie własne na podstawie programu Stories Matter

Należy jednak pamiętać, że usunięcie rozmówcy spowoduje również usunięcie wszystkich sesji i klipów przypisanych do tej osoby, a także wszystkich powiązanych z nią notatek.

5.3.2. Wprowadzanie dodatkowych informacji o rozmówcach

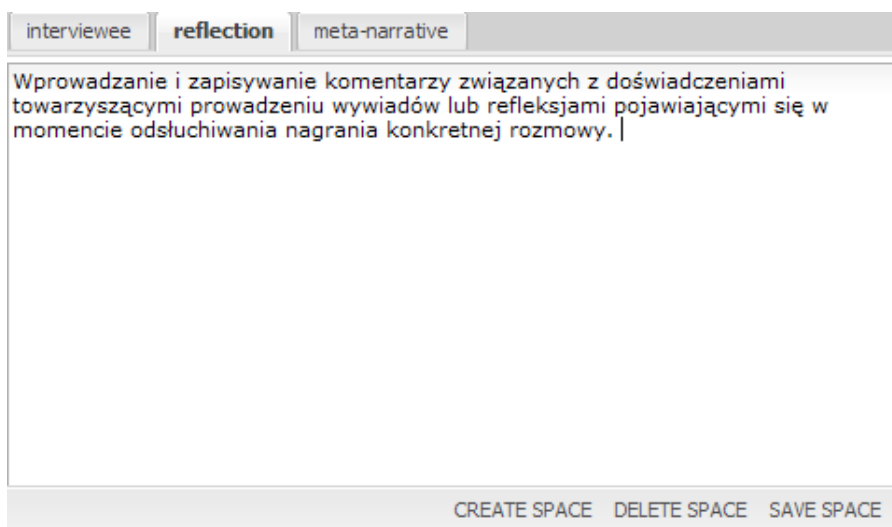
Program Stories Matter umożliwia wprowadzanie dodatkowych uzupełniających wiadomości o rozmówcach. Do dyspozycji użytkownika jest zatem formularz, w którym znajdują się trzy zakładki: *interviewee*, gdzie można zamieścić kluczowe informacje związane z kontekstem wywiadu, takie jak na przykład krótkie omówienie najważniejszych tematów podejmowanych podczas rozmowy; *biographical information* – dane biograficzne dotyczące osoby badanego czy *indexterms*, a więc wszelkie istotne terminy pozwalające na identyfikację czy przeszukiwanie rozmówców.

Ilustracja 5.9. Wygląd zakładki *interviewee*

Źródło: opracowanie własne na podstawie programu Stories Matter

W zakładce *interviewee* znajdują się dwa pola. Pierwsze to pole *summary* (podsumowanie), w którym umieszcza się krótki opis zawartości wywiadu. Użytkownicy mogą również wprowadzić znaczniki oddzielone przecinkami, zwane *comma-separated tags* (określeniami indeksowymi). Tagi te służą do dwóch celów: są one wykorzystywane podczas przeszukiwania konkretnych tematów, wydarzeń i miejsc, a także w celu generowania chmury Tagów (*tag cloud*), która będzie pomocna we wskazywaniu kluczowych tematów wewnątrz projektów oraz między nimi.

Kolejna zakładka o nazwie *reflections* pozwala na wprowadzanie i zapisywanie komentarzy związanych z doświadczeniami towarzyszącymi prowadzeniu wywiadów lub refleksjami pojawiającymi się w momencie odsłuchiwania nagrania konkretnej rozmowy.



Ilustracja 5.10. Wygląd zakładki *reflections*

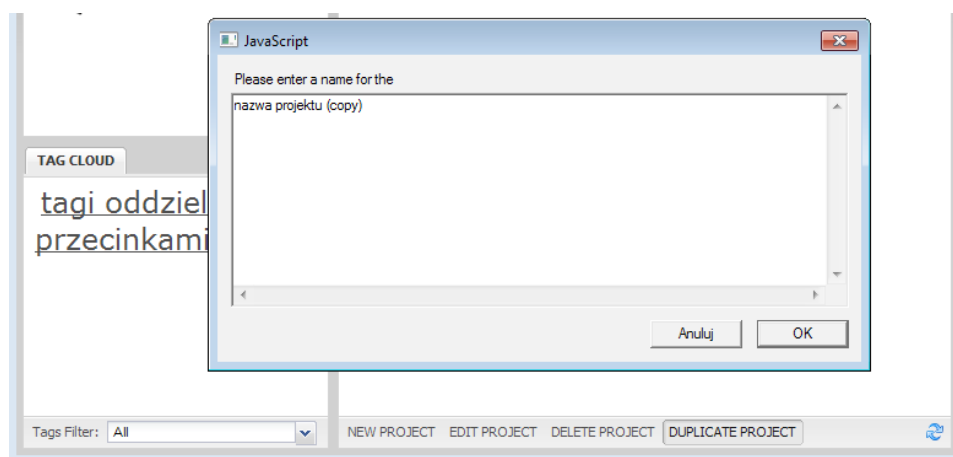
Źródło: opracowanie własne na podstawie programu Stories Matter

Trzecia z zakładek (*meta-narrative*) jest natomiast przeznaczona do wpisywania komentarzy odnoszących się do ujawnianych przez rozmówcę emocji bądź mowy ciała, które wydają się ważne z perspektywy badanych problemów.

Użytkownicy mogą również zapisywać informacje biograficzne, które udało się pozyskać podczas przeprowadzanego wywiadu. Odpowiednie do tego celu pole znajduje się w prawym górnym rogu okna programu (*interviewee biographic information*). Co więcej, program został także wyposażony w specjalne pole do zapisywania podręcznych notatek (*interviewee notes*), które zostało usytuowane w prawym dolnym rogu okna.

5.3.3. Funkcja powielania i przenoszenia wywiadów w bazie danych

Program Stories Matter wyposażony został również w dwie dodatkowe funkcje o charakterze technicznym, pozwalające dokonywać operacji na wywiadach zapisanych w bazie danych. Pierwsza z opisywanych opcji (*Duplicate Interviewee*) służy do kopiowania informacji o rozmówcy. W tym celu należy kliknąć jeden raz na daną nazwę lub obrazek jej towarzyszący, znajdujące się na liście rozmówców, i wybrać opcję duplikowania wywiadu. Użytkownik może wówczas dokonać także zmiany nazwy wywiadu lub oznaczyć go jako kopię. Po kliknięciu OK wywiad zostanie skopiowany w całości (w tym wszystkie sesje, klipy i spacje) i zapisany w liście projektów.

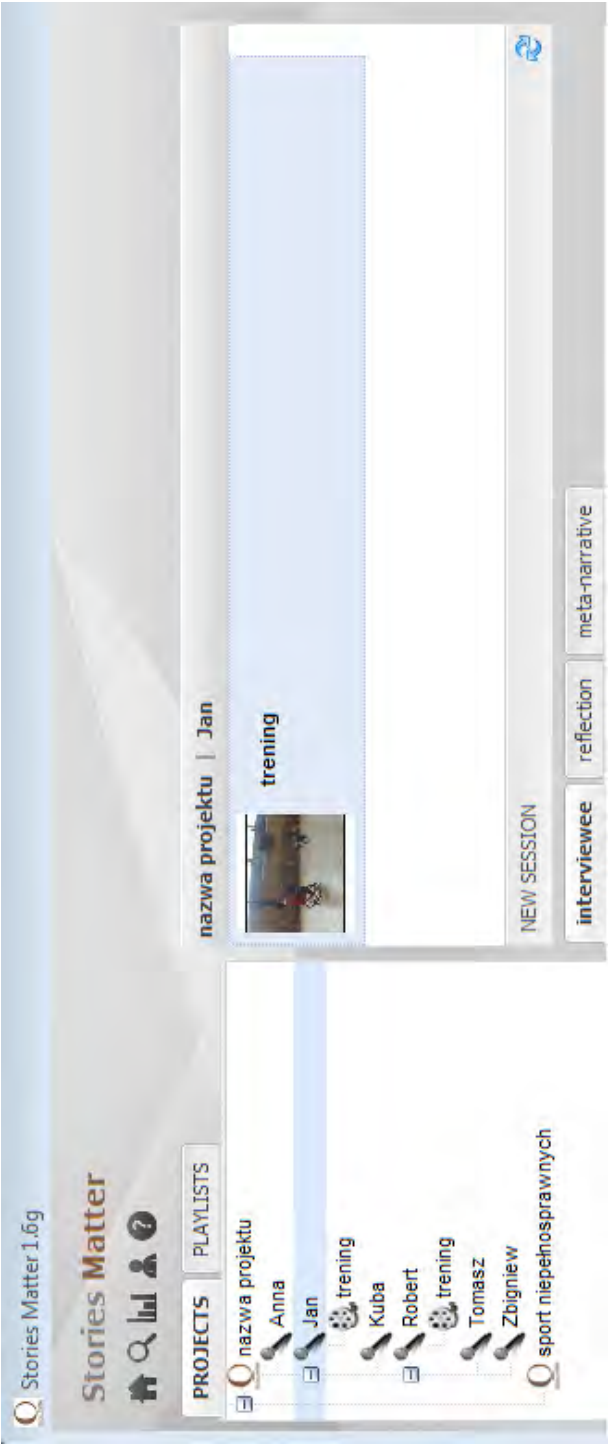


Ilustracja 5.13. Okno funkcji kopiowania wywiadu

Źródło: opracowanie własne na podstawie programu Stories Matter

Druga opcja służy do przenoszenia wywiadów w bazie danych (*Relocating Interviewees*). Możliwa jest migracja wywiadu z jednego miejsca projektu do innego w bazie danych. W tym celu należy kliknąć daną osobę z listy rozmówców, a następnie wybrać polecenie *Edit Interviewee* (Edytuj rozmowę). W dalszej kolejności należy w oknie edytora rozwinąć odpowiednie menu, wybrać projekt, do którego chce się przenieść wywiad.

Uwaga: Ta funkcja nie powieli rozmówcy. Jeśli chcesz umieścić jednego rozmówcę w więcej niż jednym projekcie, skorzystaj z funkcji *Duplicate Interviewee* i przenieś kopię wywiadu dożądanego miejsca projektu.



Ilustracja 5.14. Lista sesji dotycząca danego rozmówcy
Źródło: opracowanie własne na podstawie programu Stories Matter

5.4. Tworzenie i edycja sesji

Kolejnym poziomem („warstwą”), składającym się na wewnętrzną architekturę programu Stories Matter, są sesje (*session*), a więc poszczególne nagrania wchodzące w zakres całego wywiadu przeprowadzonego z danym rozmówcą. Lista sesji dostępna jest w środkowym (głównym) obszarze okna programu, po kliknięciu w nazwę wybranego rozmówcy.

5.4.1. Podstawowe czynności związane z dodawaniem, edytowaniem i usuwaniem sesji

Aby utworzyć nową sesję, należy wybrać przycisk znajdujący się na dole pola, w którym wyświetlana jest lista sesji prowadzonych z daną osobą. Wykonanie tej czynności spowoduje wyświetlenie okna edytora sesji, w którym użytkownik musi przypisać nazwę sesji (*Session Name*) (zazwyczaj jest to imię rozmówcy i numer sesji) oraz dodać opis sesji (*Session Description*). Następnym krokiem jest zlokalizowanie odpowiednich nośników sesji (*Session Media*), po czym należy kliknąć na przycisk *Save Session*, co spowoduje, że edytor sesji zostanie zamknięty, a plik multimedialny zacznie być importowany do programu.

Ilustracja 5.15. Okno edytora sesji

Źródło: opracowanie własne na podstawie programu Stories Matter

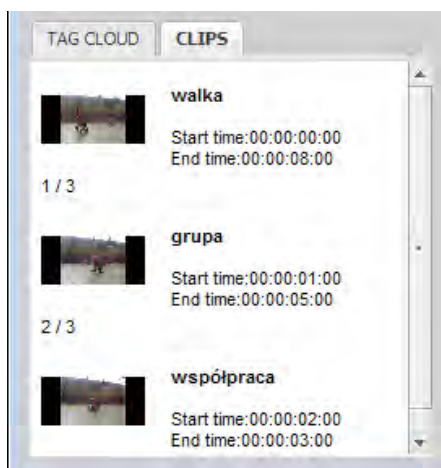
Uwaga: Importowane pliki muszą być w formacie mp3 lub FLV. Możliwe jest także importowanie plików AVI, przy czym oznacza to automatyczny proces konwersji, dokonywany podczas ładowania nowego pliku. Może to być dość długotrwały proces, przy czym użytkownik ma w jego trakcie cały czas podgląd na bieżącą sytuację poprzez widoczny pasek ładowania pliku.

Dla użytkowników komputerów Mac proces konwersji nie jest jeszcze dostępny ze względu na ograniczenia w programie Adobe Air. Przy czym w tej sytuacji można samodzielnie wykonać konwersję pliku za pomocą jednego z wielu programów *open source*, jakie można znaleźć w Internecie.

Należy też pamiętać, aby wszystkie konwertowane pliki były mniejsze niż 600 MB i zawierały co najmniej trzy klatki na sekundę, dzięki czemu powstałe pliki FLV będą wystarczająco dobrej jakości, aby zapewnić optymalną wydajność w programie.

5.4.2. Przeglądanie i wprowadzanie dodatkowych informacji o sesjach

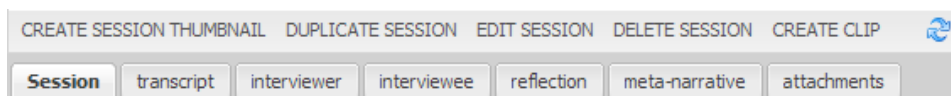
Program Stories Matter umożliwia sprawne i łatwe w obsłudze przeglądanie zawartości sesji. Wystarczy kliknąć na wybraną nazwę sesji bądź jej obrazek znajdujące się na liście sesji, co spowoduje wyświetlenie listy klipów (*clip list*) tej sesji w lewym dolnym polu okna programu, w którym standardowo znajduje się chmura tagów (*tag cloud*).



Ilustracja 5.16. Okno programu z wyświetloną listą klipów wybranej sesji

Źródło: opracowanie własne na podstawie programu Stories Matter

Jednokrotne kliknięcie na ikonę sesji lub jej opis spowoduje również dostęp do opcji edycji i usuwania sesji. Przy czym należy pamiętać, że usunięcie sesji spowoduje również usunięcie wszystkich plików oraz notatek z nią powiązanych. Z kolei dwukrotne kliknięcie na ikonę sesji lub jej opis umożliwia użytkownikowi rozpoczęcie dodawania informacji związanych z daną sesją. Wówczas też zostanie wyświetlone okno z siedzioma widocznymi zakładkami, w których można wprowadzić dodatkowe informacje dotyczące sesji.



Ilustracja 5.17. Pasek z dostępnymi opcjami edytowania sesji

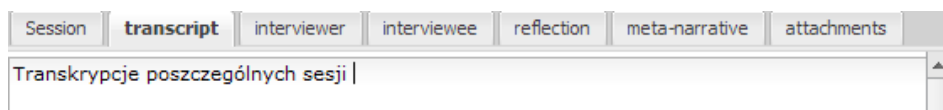
Źródło: opracowanie własne na podstawie programu Stories Matter

Pierwsza zakładka (*Session*) zawiera pola pozwalające na zapisanie podsumowujących informacji o kontekście i warunkach przeprowadzenia sesji (*Summary*). Ponadto istnieje tutaj możliwość określenia lokalizacji (*Location*) oraz wyświetlania miejsca, w którym została zarejestrowana sesja za pomocą *Google maps*. W tym celu po wprowadzeniu dokładnej nazwy lokalizacji w przeznaczone do tego pole należy kliknąć na przycisk znajdujący się obok, co spowoduje wyświetlenie dodatkowego okna, zawierającego mapy dostępne za pośrednictwem przeglądarki Google. W zakładce *Session* można określić datę, a także język sesji oraz format zapisu danego nagrania czy listę tagów.

Ilustracja 5.18. Wygląd oraz pola zakładki *Session*

Źródło: opracowanie własne na podstawie programu Stories Matter

Kolejna zakładka *transcript* umożliwia użytkownikowi włączenie transkrypcji poszczególnych sesji wywiadów.



Ilustracja 5.19. Wygląd okna *Transcript*

Źródło: opracowanie własne na podstawie programu Stories Matter

Trzecia w kolejności zakładka *interviewer* daje użytkownikom możliwość zachowania informacji o osobie przeprowadzającej wywiad, w tym takie dane, jak imię i nazwisko, adres, datę urodzenia i dane kontaktowe.

Ilustracja 5.20. Wygląd okna *interviewer*

Źródło: opracowanie własne na podstawie programu Stories Matter

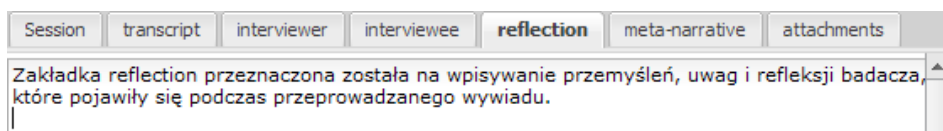
Z kolei „bliźniacza” zakładka *interviewee* przeznaczona jest do wpisywania tożsamyh informacji, ale o osobie, z którą wywiad został przeprowadzony.

Ilustracja 5.21. Wygląd okna *interviewee*

Źródło: opracowanie własne na podstawie programu Stories Matter

Trzeba przy tym pamiętać, że dane wprowadzone zarówno do jednej, jak i do drugiej zakładki wymagają szczególnej ochrony, co zapewnia sam program, bowiem dane te nie są dostępne innym użytkownikom podczas pracy on-line.

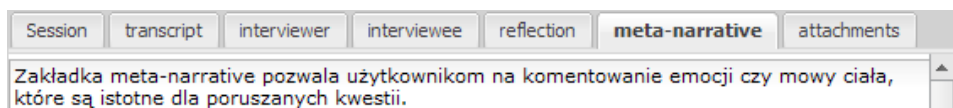
Następna zakładka *reflection* przeznaczona została na wpisywanie przemyśleń, uwag i refleksji badacza, które pojawiły się podczas przeprowadzanego wywiadu.



Ilustracja 5.22. Wygląd okna *reflection*

Źródło: opracowanie własne na podstawie programu Stories Matter

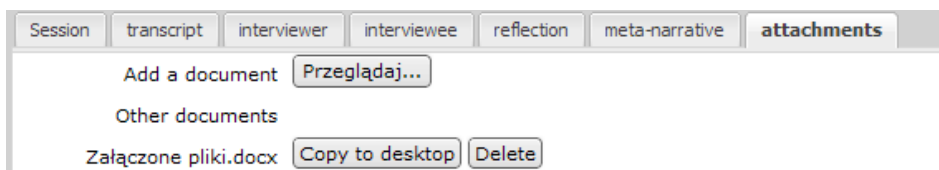
Podobną funkcję do powyższej zakładki posiada zakładka *meta-narrative*, która pozwala użytkownikom na komentowanie emocji czy mowy ciała, które są istotne dla poruszanych kwestii.



Ilustracja 5.23. Wygląd okna *meta-narrative*

Źródło: opracowanie własne na podstawie programu Stories Matter

Wreszcie ostatnią zakładką jest *attachments*, a więc przestrzeń, do której można importować zdjęcia, filmy, a także wszelkiego rodzaju materiały i dokumenty, które związane są z sesją przeprowadzonego wywiadu.



Ilustracja 5.24. Wygląd okna *attachments*

Źródło: opracowanie własne na podstawie programu Stories Matter

Aby dodać załącznik, należy wybrać klawisz Przeglądaj (*browse*). Wówczas też pojawi się odpowiednie okienko, w którym użytkownik będzie w stanie dokonać wyboru polegającego na dołączeniu dokumentu (pliku). Ponadto program daje możliwość określenia przez użytkownika, jakie zakładki pojawią się w oknie

służącym do opisu sesji. W tym celu należy kliknąć na *create space*, a następnie w wyświetlonym oknie (*Space Editor*) wybrać konkretne opcje zakładek spośród tych, które znajdują się w rozwijanej liście.

Użytkownik programu może dodać notatkę dotyczącą konkretnej sesji, która pojawi się w polu widocznym w prawym dolnym rogu programu (*session notes*). W taki sam sposób można wprowadzić informację o historii życia rozmawiającego, która zostanie zamieszczona w górnym prawym polu okna programu (*interviewee biographical information*).

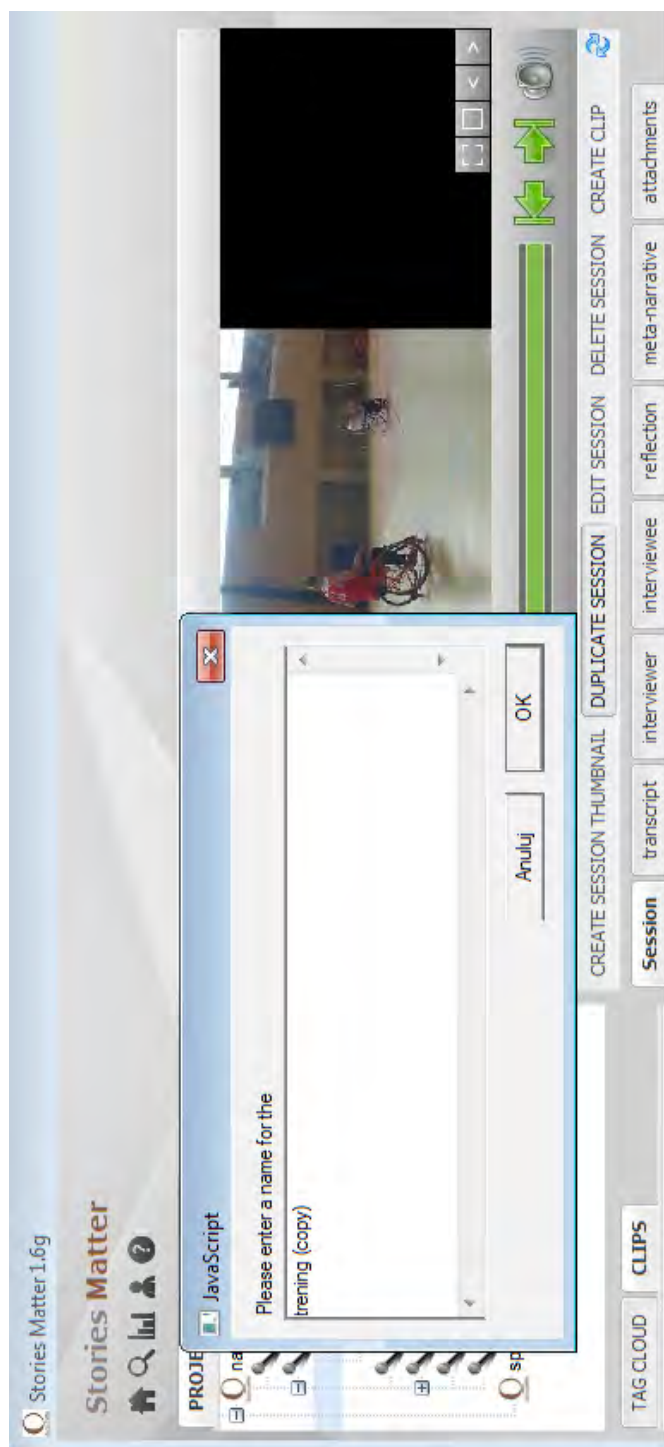
The image shows a vertical window with two main sections. The top section is titled 'Interviewee biographical information' and contains a text area with the placeholder text 'Informacje o historii życia rozmawiającego.' Below this text area is a 'SAVE' button. The bottom section is titled 'Session notes' and contains a text area with the placeholder text 'Notatki dotyczącą konkretnej sesji.' Below this text area is a 'SAVE NOTES' button.

Ilustracja 5.25. Pola *session notes* i *interviewee biographical information*

Źródło: opracowanie własne na podstawie programu Stories Matter

5.4.3. Funkcja powielania i przenoszenia wywiadów w bazie danych

Podobnie jak przy omawianych wcześniej rozmówcach, także w przypadku sesji program Stories Matter wyposażony został również w dwie dodatkowe funkcje o charakterze technicznym. Pierwsza z opisywanych opcji (*Duplicate Session*) służy do kopiowania informacji o sesji, a jej zastosowanie polega na kliknięciu danej nazwy sesji spośród tych, które znajdują się na wyświetlanej liście, następnie wybraniu opcji kopiowania (*Duplicate*) i określeniu nazwy sesji w odpowiednim oknie. Po kliknięciu OK sesja zostanie skopiowana w całości (wraz ze wszystkimi nagraniami i zakładkami z wprowadzonymi uprzednio przez użytkownika informacjami o sesji) oraz zapisana w liście projektów.



Ilustracja 5.26. Okno funkcji kopiowania sesji

Źródło: opracowanie własne na podstawie programu Stories Matter

Druga opcja służy do przenoszenia sesji w bazie danych (*Relocating Session*). Możliwa jest migracja sesji z jednego miejsca projektu do innego w bazie danych. W tym celu należy kliknąć daną nazwę sesji z listy, a następnie wybrać polecenie Edytuj (*Edit Session*). W dalszej kolejności należy w oknie edytora rozwinąć odpowiednie menu, wybrać projekt, w którym chce się przenieść sesję.

Uwaga: Ta funkcja nie powiela sesji, dlatego jeśli chcemy umieścić jedną sesję w więcej niż jednym projekcie, musimy skorzystać z funkcji *Duplicate Session* i przenieść kopię sesji dożądanego miejsca projektu.

5.5. Tworzenie, edycja i eksportowanie klipów

Ostatnim poziomem („warstwą”), składającym się na wewnętrzną architekturę programu Stories Matter są klipy (*clips*), które mają umożliwić użytkownikom zlokalizowanie konkretnych punktów w sesji, w której rozmawiający omawia temat, pasujący do konkretnych zainteresowań badawczych użytkownika.

5.5.1. Podstawowe czynności związane z tworzeniem, edytowaniem i eksportowaniem klipów

Użytkownik może utworzyć klip, wybierając przyciski *clip in* oraz *clip out*, które pozwalają na określenie fragmentu nagrania, jaki ma zostać „wydzielony” z sesji. Odbyna się to w ten sposób, że podczas odtwarzania sesji należy najpierw kliknąć na przycisk *clip in*, co wskazuje na punkt początkowy oznaczenia, a następnie na przycisk *clip out*, który zamyka wybrany fragment sesji.



Ilustracja 5.27. Przyciski *clip in* oraz *clip out*

Źródło: opracowanie własne na podstawie programu Stories Matter

Następnie należy wybrać opcję tworzenia klipu (*create clip*) z opcji znajdujących się na dole odtwarzacza multimedialnego, co spowoduje wyświetlenie odpowiedniego okna programu, w którym użytkownik będzie mógł wprowadzić nazwę klipu (*clip name*), jego opis (*clip description*) oraz nadać wyróżniające go tagi (*comma separated tags*).

Ilustracja 5.28. Okno tworzenia klipów

Źródło: opracowanie własne na podstawie programu Stories Matter

Istniejące klipy mogą być edytowane przez wybranie funkcji edycji klipu lub usunięcie klipu z opcji znajdujących się na dole odtwarzacza multimedialnego. Użytkownicy mogą również wygenerować miniaturę klipu, wybierając odpowiednią do tego celu opcję (*create clip thumbnail*). Miniatura pojawi się w obszarze listy klipów, towarzysząc nazwie danego klipu.



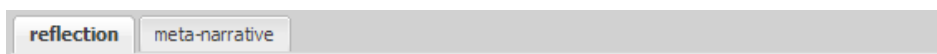
Ilustracja 5.29. Pasek opcji edycji klipów

Źródło: opracowanie własne na podstawie programu Stories Matter

Użytkownicy mają również możliwość eksportowania klipu lub ich listy na pulpit komputera, co ułatwia integrację z prezentacjami programu PowerPoint. W tym celu użytkownicy powinni kliknąć na żądany klip z listy klipów, a następnie wybrać opcję wyeksportuj klip (*export clip*) z opcji wyświetlanych poniżej odtwarzacza multimedialnego. Na pulpicie zostanie utworzony folder o nazwie Stories Matter, a w nim będzie to klip w formacie .avi.

Dodatkową cechą Stories Matter jest lista klipów, którą można wyświetlić w lewym dolnym rogu programu. Wszelkie klipy utworzone podczas sesji zostaną wyświetlone na liście klipów znajdującej się w polu po lewej stronie odtwarzacza multimedialnego. Na liście klipów widoczna będzie miniatura, nazwa i opis każdego klipu. Dwukrotne kliknięcie konkretnego klipu na liście spowoduje jego załadowanie w odtwarzaczu multimedialnym, co także umożliwi użytkownikowi edycję lub usunięcie klipu.

Użytkownicy mogą również wprowadzić dodatkowe informacje o klipie w przeznaczone do tego celu pola, a więc znane już z poprzednich opcji: *reflection* czy *meta-narrative*.



Ilustracja 5.30. Okno z polami przeznaczonymi do wprowadzenia dodatkowych informacji o klipie

Źródło: opracowanie własne na podstawie programu Stories Matter

Przy czym tagi oddzielone przecinkami należy wprowadzać bezpośrednio w przestrzeni refleksji. Warto je wprowadzić, bowiem program będzie na ich podstawie przeszukiwał dane, a także generował chmurę tagów w celu podkreślenia kluczowych tematów narracji.

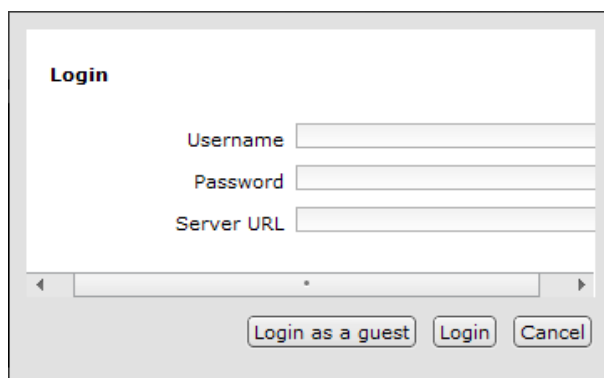
Użytkownicy mogą wpisywać swoje notatki o klipie (*notes*), które pojawiają się w prawym dolnym rogu okna programu, a także dokonywać operacji polegających na przenoszeniu klipów w ramach projektu, co odbywa się w sposób tożsamy do tego, jaki był opisywany przy prezentowaniu opcji dostępnych podczas edycji sesji oraz rozmówców.

5.6. Funkcje on-line dostępne w programie Stories Matter

Cechą Stories Matter jest to, że wielu użytkowników może współpracować nad projektem zdalnie poprzez zdalny dostęp i łączenie baz danych. Przy czym, aby rozpocząć współpracę, należy posiadać dostęp do serwera i zainstalować na nim oprogramowanie Stories Matter. W tym celu zaleca się skorzystanie z fachowej pomocy informatyka zaznajomionego ze środowiskiem Adobe Air, aby odpowiednio skonfigurować wszystkie niezbędne parametry. Aby uzyskać dostęp do serwera Stories Matter, konieczne będzie także przypisanie administratora oraz nadawanie różnych ról użytkownikom, takie jak Menedżer projektów, wydawca i goście. W ten sposób tylko zarejestrowani użytkownicy będą mieli dostęp do wybranych projektów. Dla użytkowników, którzy chcą pracować samodzielnie, te dodatkowe kroki nie są konieczne, ponieważ oprogramowanie może być nadal używane w trybie off-line.

5.6.1. Uzyskiwanie dostępu on-line

Użytkownicy mogą rozpocząć pracę on-line, wybierając login z prawego górnego rogu okna programu, co spowoduje pojawienie się okna, w którym należy wpisać nazwę użytkownika, hasło i adres URL serwera. Użytkownicy mają również możliwość wyboru loginu jako gościa, co umożliwi im dostęp do danego serwera bez konieczności roli użytkownika i hasła.



Ilustracja 5.31. Okno logowania

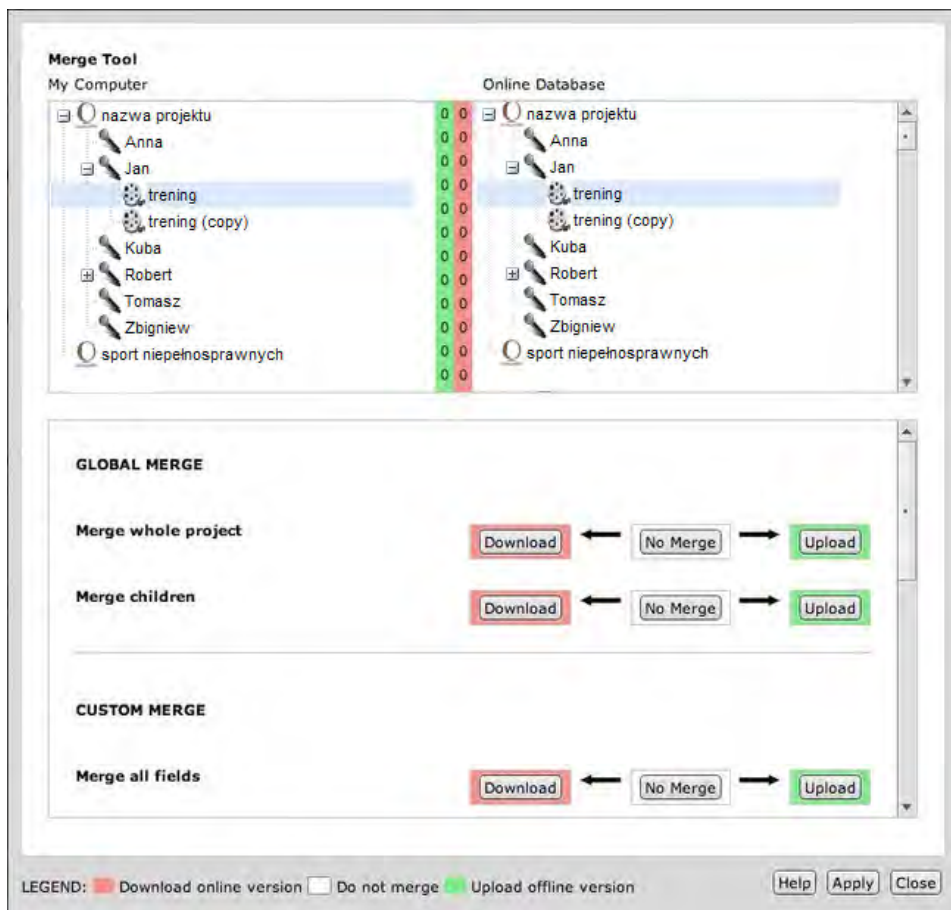
Źródło: opracowanie własne na podstawie programu Stories Matter

Po zalogowaniu użytkownicy będą mogli przeglądać i współdziałać z bazą danych na serwerze hosta, a baza danych zbudowana na domowym komputerze nie będzie już widoczna. Użytkownicy mogą w dowolnym momencie wrócić do lokalnej bazy danych, wybierając *logout*.

5.6.2. Narzędzie scalania bazy danych

Podczas pracy on-line użytkownicy będą mogli połączyć elementy lokalnej bazy danych z bazą danych umieszczoną na serwerze. Aby uzyskać dostęp do narzędzia scalania, należy wybrać ikonę wykresu w lewym górnym rogu okna programu. Użytkownicy uzyskają także dostęp do serii narzędzi baz danych, gdzie będą między innymi mogli wyświetlać opcje eksportowania i importowania plików. Aby rozpocząć łączenie informacji między serwerem a komputerem lokalnym, należy wybrać narzędzie scalania, co spowoduje wyświetlenie się odpowiedniego okna programu.

W oknie narzędzia scalania po lewej stronie znajduje się lista projektów, które dostępne są w bazie danych istniejącej na lokalnym komputerze, natomiast lista projektów widoczna po prawej stronie pokazuje bazę danych zapisaną na



Ilustracja 5.32. Okno narzędzia Merge Tools

Źródło: opracowanie własne na podstawie programu Stories Matter

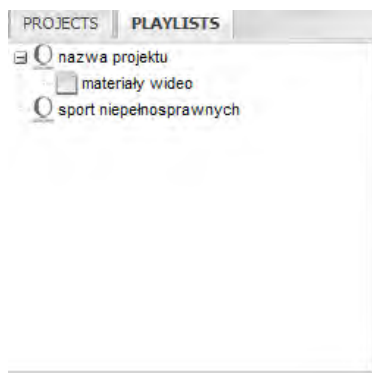
serwerze internetowym. Kolorem szarym oznaczono te wszystkie elementy, które nie zostały jeszcze połączone i z tego powodu są niedostępne. Podstawowe instrukcje są wymienione poniżej listy projektów i można je wyświetlić w dowolnym momencie, wybierając opcję pomocy. Aby połączyć konkretne elementy baz, należy wybrać określoną pozycję zapisaną kolorem czarnym. Spowoduje to wyświetlenie się okna programu z opcjami służącymi do scalania. Jeśli zaś zechcemy przesłać projekt, wywiad lub sesję z lokalnego komputera do bazy danych on-line, należy wybrać polecenie **Prześlij**. Aby pobrać projekt, wywiad lub sesję z bazy danych on-line do lokalnego komputera, wybierz opcję **Pobierz**. Użytkownicy mają możliwość łączenia całych projektów, jak również scalania lub łączenia konkretnych elementów.

5.7. Dodatkowe funkcje programu Stories Matter

Program Stories Matter został wyposażony w kilka dodatkowych funkcji, co zdecydowanie podniosło jego walory użytkowe. Należą do nich: możliwość tworzenia i edytowania list odtwarzania, generowania chmur tagów, eksportowania klipów do programu PowerPoint lub tworzenia kopii zapasowych zawartości bazy danych w dokumencie html, aż wreszcie przeszukiwania bazy danych przy użyciu kluczowych haseł.

5.7.1. Listy odtwarzania

Pierwszą funkcją jest możliwość tworzenia list odtwarzania. Użytkownicy mogą tworzyć playlisty z klipów, wybierając przeglądarkę playlist z opcji znajdujących się u lewym górnym rogu okna programu. Następnie należy wybrać projekt, w ramach którego zostanie utworzona lista odtwarzania spośród istniejących klipów.

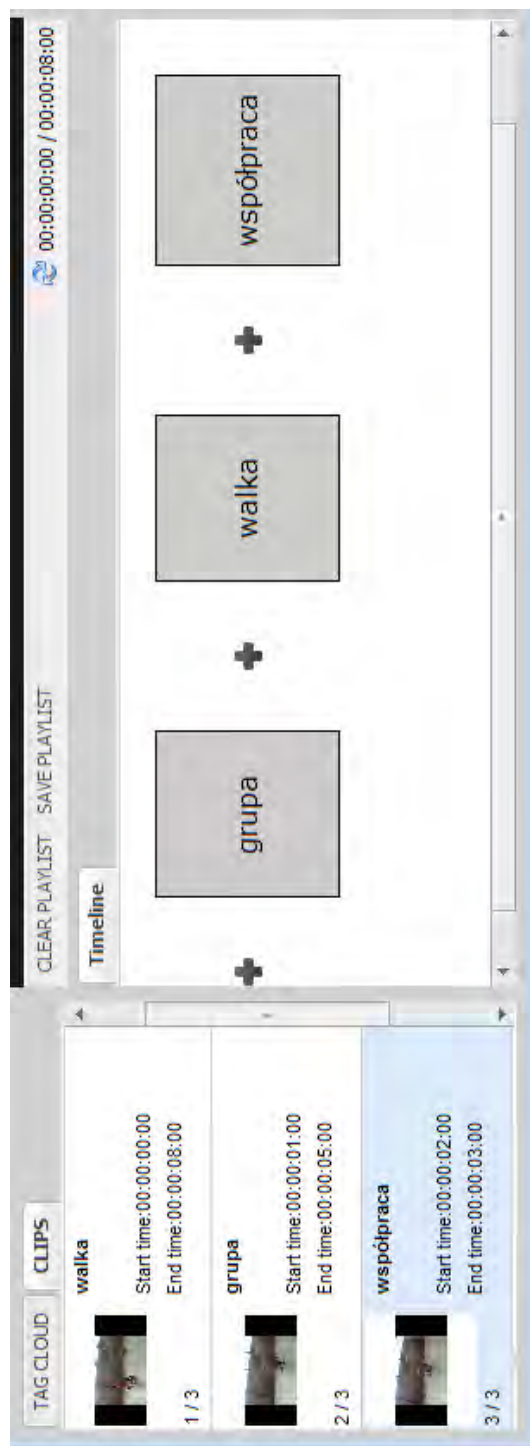


Ilustracja 5.33. Lista odtwarzania klipów (playlista)

Źródło: opracowanie własne na podstawie programu Stories Matter

Natomiast wybierając danego rozmówcę, użytkownicy będą mogli wyświetlić pełną listę klipów dla każdej sesji. Aby utworzyć listę odtwarzania, należy wybrać „+” na osi czasu. Za każdym razem, gdy użytkownik wybiera znaczek plusa, na osi czasu pojawi się pusty kwadrat. Teraz, aby dodać w miejsce wspomnianego pustego kwadratu dany klip, należy kliknąć na jego nazwę. Postępując każdorazowo w ten sam sposób, użytkownik może stworzyć własną listę odtwarzania klipów.

Następnym krokiem jest zapisanie playlisty, co następuje po kliknięciu odpowiedniej opcji zapisu znajdującej się na dole odtwarzacza multimedialnego. Użytkownicy mogą także usunąć listę odtwarzania, wybierając opcję wyczyść.

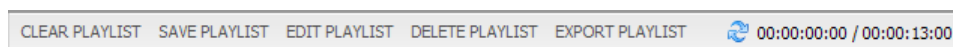


Ilustracja 5.34. Dodawanie klipów do playlisty

Źródło: opracowanie własne na podstawie programu Stories Matter

Po kliknięciu na opcję *Save* otworzy się okno edytora listy odtwarzania, w którym należy wpisać nazwę i podać jej krótki opis.

Po tych czynnościach nowo utworzona lista zostanie wyświetlona w przeglądarce playlist. Kliknięcie raz na istniejącej liście odtwarzania widocznej w przeglądarce playlist spowoduje, że pojawią się dodatkowe opcje w dolnej części osi czasu. Korzystając z tych opcji, użytkownicy będą mogli m.in. zapisywać zmiany, a także edytować lub usunąć listę odtwarzania.



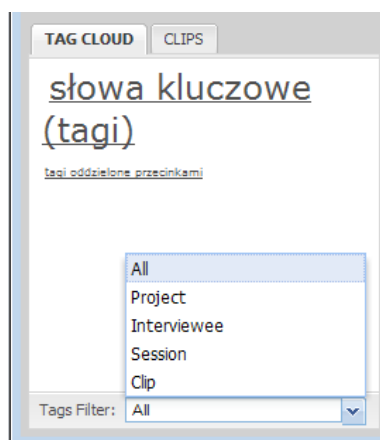
Ilustracja 5.35. Dodatkowe opcje zapisywania, edycji lub usuwania playlisty

Źródło: opracowanie własne na podstawie programu Stories Matter

Aby odtworzyć nagrania zapisane w playliście, należy kliknąć na przycisk odtwarzania w odtwarzaczu multimedialnym. Klipy będą odtwarzane jeden za drugim, do momentu, aż użytkownik nie wybierze pauzy. Dodatkowo do danej listy można stworzyć jej opis (*playlist description*) oraz wprowadzić notatki dotyczące całej listy bądź poszczególnych, składających się na nią klipów (*playlist notes*).

5.7.2. Chmura Tagów

Za każdym razem, gdy użytkownicy dodają tagi oddzielone przecinkami, dotyczące poszczególnych rozmówców, sesji czy klipów, oprogramowanie generuje chmurę tagów.



Ilustracja 5.36. Chmura Tagów (znaczników)

Źródło: opracowanie własne na podstawie programu Stories Matter

Jeśli użytkownik będzie chciał sprawdzić, na jakich kwestiach koncentrują się rozmówcy, może posłużyć się wprowadzonymi uprzednio znacznikami. Im większą czcionką będą one zapisane, tym częściej będą pojawiać się w wypowiedzi rozmówcy. I odwrotnie, te rzadziej obecne w wypowiedzi danego rozmówcy pojawiają się jako wyrażenia zapisane małą czcionką.

Użytkownicy mogą wybrać wyświetlanie chmury tagów dla danego projektu, konkretnego rozmówcy lub sesji, wybierając żądaną opcję filtrowania tagów, znajdującą się u podstawy chmury tagów.

5.7.3. Dodatkowe narzędzia bazy danych

Poza opisanym już wcześniej narzędziem scalania dostępne są dodatkowe narzędzia bazy danych, gdy użytkownik wybierze ikonę wykresu z menu.

Po pierwsze użytkownicy mają możliwość eksportowania zawartości swojej bazy danych w formacie HTML, co w praktyce umożliwia ich łatwą integrację z witryną internetową. Aby rozpocząć proces, należy wybierać eksport do HTML. Następnie użytkownicy zostaną poproszeni o wybranie lokalizacji dokumentów zapisanych w formacie HTML.

Uwaga: Ponieważ opcja eksportu zawartości bazy danych w formacie HTML tworzy cztery oddzielne foldery i indeks, warto je zapisać razem w specjalnie do tego celu utworzonym nowym folderze na dysku komputera.

Druga opcja eksportowania zawartości lokalnej bazy danych służy do tworzenia kopii zapasowej. Natomiast trzecia daje możliwość przywrócenia bazy danych do określonego stanu poprzez wykorzystanie opcji importowania kopii zapasowej i zastąpienia jej istniejącą bazą danych.

5.7.4. Opcje przeszukiwania

Program Stories Matter został także wyposażony w opcję wyszukiwania konkretnych typów treści. Wybór ikony szkła powiększającego z menu w lewym górnym rogu okna oprogramowania spowoduje wyświetlenie nowego okna, w którym można prowadzić dokładne lub ogólne wyszukiwanie w ramach istniejącej bazy danych.

Użytkownicy mogą przeszukiwać bazę danych, wprowadzając hasła w polu zatytułowanym wyszukiwane hasła. Korzystając z opcji przeszukiwania, można także określić jego specyfikację poprzez wybranie odpowiednich opcji przeszukiwania.

Here you have various tools to manage your database:

Export to HTML

Create an HTML page with your local database, in the directory you select.

[Export to HTML](#)

Export files

Create a backup of your local database in a specified directory. You can restore a database using the tool below (Import Database).

[Export files](#)

Import Files

Restore a backup of a database from a specified directory. You can backup your local database using the tool above (Export Database).

[Import Files](#)

WARNING: This will overwrite the current files

Ilustracja 5.37. Narzędzia bazy danych

Źródło: opracowanie własne na podstawie programu Stories Matter



Ilustracja 5.38. Funkcja wyszukiwania w programie Stories Matter
Źródło: opracowanie własne na podstawie programu Stories Matter

Wyniki wyszukiwania zostaną uporządkowane według kolejności: rozmówcy, sesje lub klipy. Kliknięcie żądanego wyniku spowoduje przeniesienie użytkownika bezpośrednio do odpowiedniego okna i umożliwi mu natychmiastowe odsłuchanie danego klipu lub sesji. Warto zauważyć, że drugi wygodny sposób na dostęp do wyszukiwarki to po prostu kliknięcie określonego tagu wyświetlanego w chmurze tagów.

5.8. Uwagi końcowe

Program Stories Matter to narzędzie, które może pomóc w procesie porządkowania i systematyzowania zebranych materiałów. Dzięki systemowi notatek oraz opisów pozwala na wprowadzanie dodatkowych informacji o gromadzonych plikach i ich zawartości. Natomiast narzędzie przeszukiwania oraz tworzenia chmury tagów zdecydowanie usprawnia proces rozpoznawania kluczowych tematów i kwestii poruszanych przez rozmówców.

Niewątpliwym atutem programu jest jego intuicyjność i prostota użycia, co pozwala niemalże każdemu, nawet niewprawnemu badaczowi na szybkie zaznajomienie się z możliwościami oprogramowania i szybkie jego zastosowanie, bez konieczności długotrwałego procesu uczenia się obsługi. Pewnym mankamentem może być natomiast ograniczony zakres formatów plików, jakie mogą być importowane do programu, a także konieczność dość pracochłonnego ich opisywania, zwłaszcza gdy chce się wprowadzać wszystkie dane oraz informacje o każdym z rozmówców czy specyfice wywiadu i jego przebiegu.

Niemniej jednak jest to program, który z całą pewnością przyda się wszystkim tym, którzy myślą o realizacji badań z zastosowaniem danych audiowizualnych oraz stosujących w swojej pracy badawczej metodę biograficzną lub po prostu opierających się na historiach życia.

Program Stories Matter został stworzony i udostępniony przez Centre for Oral History and Digital Storytelling.

Zarówno sam program, jak i dodatkowe informacje o nim można pobrać ze strony:

<http://storytelling.concordia.ca/development/>

Rozmiar pliku 37,2 MB (wersja Stories Matter 1.1)

Przed instalacją programu należy upewnić się, czy na komputerze posiadamy oprogramowanie Adobe Air 1.1. Jeśli nie, można je pobrać ze strony: <http://get.adobe.com/air/> (dla systemów linux: <http://labs.adobe.com/technologies/air/>)

Systemy, pod którymi można uruchomić program: Windows (począwszy od wersji XP), Mac OS X (od wersji v.10.4.9) oraz Linux (od wersji Fedora Core 8, Ubuntu 7.10, Open Suse 10.3)

Przy opracowywaniu materiału o programie posłużono się instrukcją użytkownika autorstwa Jacques Langlois z Centre for Oral History and Digital Storytelling Concordia University.

6. ELAN – oprogramowanie do analizy jakościowej materiałów audiowizualnych

ELAN to narzędzie do analizy jakościowej materiałów audiowizualnych, które pozwala tworzyć, edytować, wizualizować i wyszukiwać adnotacje w przypadku danych wideo i audio. Program został opracowany w Instytucie Psycholingwistyki Maxa Plancka, Nijmegen w Holandii w celu zapewnienia solidnych podstaw technologicznych do opisu i wykorzystania nagrań multimedialnych. ELAN jest przeznaczony specjalnie do analizy języka, w tym także do języka migowego i gestów, ale może być używany przez wszystkich, którzy pracują z danymi wideo i/lub audio w celach analizy oraz dokumentacji danych.

Proces opracowywania i analizy danych obejmuje trzy kluczowe etapy: definiowanie typów warstw i poziomów, wybieranie odstępów czasowych oraz wprowadzanie adnotacji. Dlatego ten krótki przewodnik jest zorganizowany wokół tych trzech kroków w taki sposób, aby pomóc zrozumieć logikę działania programu oraz sposób używania jego głównych funkcji.

6.1. Kluczowe elementy projektu i ich definicje

6.1.1. Adnotacje

Pierwszym kluczowym elementem organizującym pracę użytkownika programu są adnotacje, a więc dowolne typy tekstu (np. transkrypcja, tłumaczenie, kodowanie itp. Jest ona przypisywana wybranemu odstępowi czasu pliku wideo/audio (np. wypowiedzeniu rozmówcy) lub adnotacji na innym poziomie (np. tłumaczenie jest przypisywane do transkrypcji ortograficznej).

6.1.2. Poziomy

Poziomy (tier) to zbiory adnotacji mających te same cechy, na przykład jeden poziom może zawierać transkrypcję ortograficzną, a inny tłumaczenie. Poziom może być bezpośrednio związany z przedziałem czasowym pliku multimedialnego lub może „odwoływać się” do innego poziomu i w ten sposób umożliwiać budowanie ich hierarchii.

6.1.3. Rodzaje (typy) poziomów

Poziomy podzielone są według typów, które określają charakter oraz sposób organizacji adnotacji. Mogą występować następujące przypadki: brak (a więc sytuacja, w której występują wolne i niezależne adnotacje), podział czasu (adnotacje są powiązane z osią czasu), podział symboliczny (adnotacje są podzielne według odwołań, ale nie mogą być powiązane z osią czasu), stowarzyszenie symboliczne (jedna adnotacja na danym poziomie odpowiada dokładnie jednej adnotacji na innym poziomie).

6.1.4. Szablony

Szablon to specjalny plik do pisania adnotacji, zawierający zestawienie poziomów służących do robienia adnotacji w określony sposób. Oszczędza to dużo pracy, aby opatrzyć nowy plik adnotacji na szablonie, i ułatwia opisanie kilku filmów w ten sam sposób.

6.1.5. Plik adnotacji (*.eaf)

Plik adnotacji jest dokumentem, który zawiera wszystkie informacje o poziomach (ich atrybuty i relacje między nimi), a także adnotacje oraz wyrównanie czasu i linki do plików multimedialnych.

6.1.6. Plik multimedialny (*.mpg, *.wav, itd.)

Plik multimedialny zawiera zdigitalizowany plik wideo/audio (np. *.mpg) lub tylko dane audio (*.wav), które chce się skomentować. Posiada oś czasu, z którą połączony jest plik adnotacji (*.eaf).

6.2. Wewnętrzna organizacja plików w programie

Każdy projekt ELAN składa się z co najmniej dwóch plików: jednego (lub więcej) pliku multimedialnego i jednego pliku, w którym umieszczane są adnotacje. Mam zatem jeden lub więcej plików multimedialnych: wideo (*.mpg itp.) i/lub audio (*.wav). Plik wideo umożliwia oglądanie filmu i odsłuchiwanie dźwięku. Przy czym, aby zobaczyć przebieg fali, trzeba utworzyć dodatkowy plik *.wav za pomocą programu konwersji, który konwertuje dane audio z pliku *.mp(e)g na format *.wav. W przypadku wielu plików wideo odtwarzany jest dźwięk

pierwszego z nich. Rodzaj i liczba obsługiwanych formatów wideo zależy od ustawień komputera i zainstalowanych kodeków. Natomiast pojedynczy plik adnotacji utworzony zostaje bezpośrednio w programie (z rozszerzeniem *.eaf, „EUDICO Annotation Format”) lub może być do niego importowany. To właśnie w tym pliku są zapisywane wszystkie informacje – nigdy zaś w pliku multimedialnym. Dlatego dodanie adnotacji w programie ELAN nie zmienia w żaden sposób pliku multimedialnego.

6.3. Okno programu ELAN

W oknie programu na samej górze widoczny jest pasek głównego menu z pozycjami: *File, Edit, Annotation, Tier, Type, Search, View, Options, Window, Help*. Poniżej wspomnianego menu znajduje się obszar wyświetlania materiałów filmowych, pod nim cały zestaw przycisków służących do kontroli przeglądania danych.

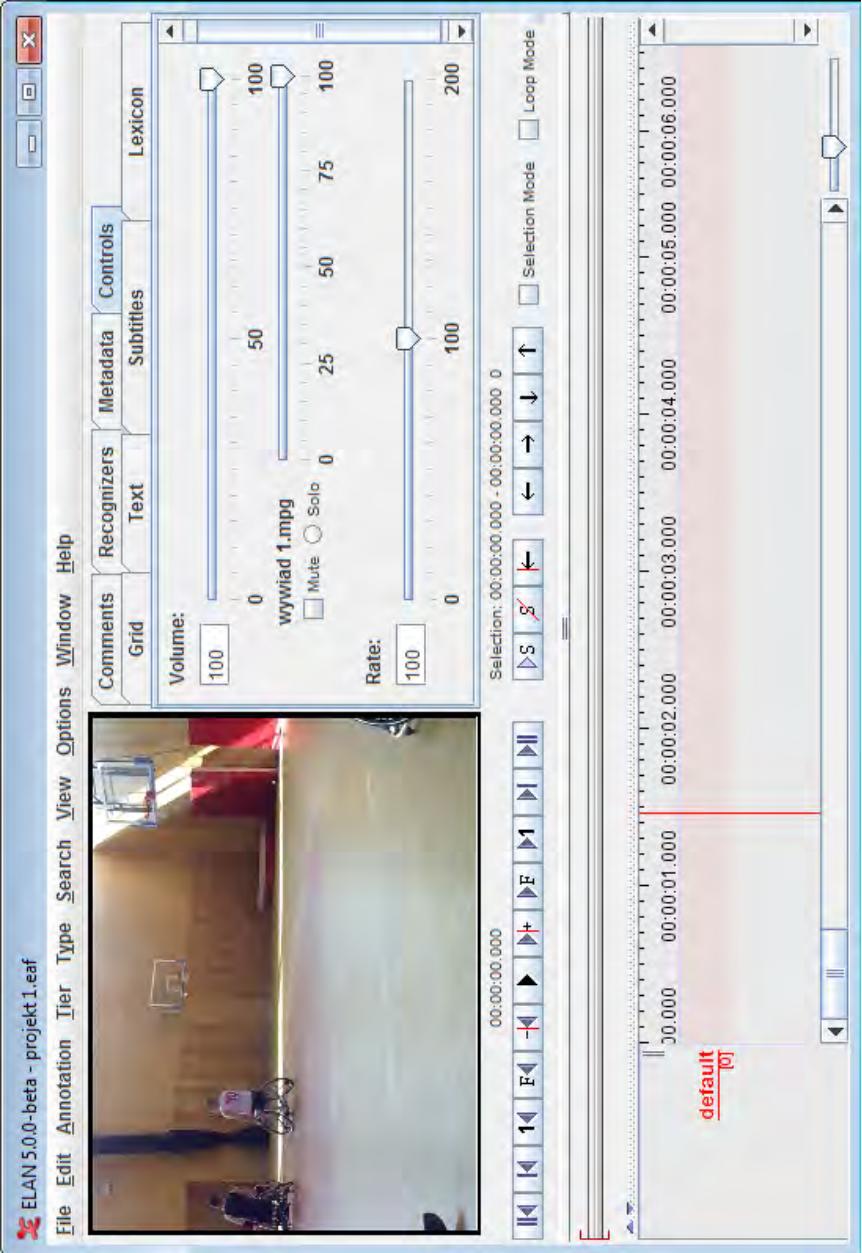
W środkowej części okna znajduje się pasek pokazujący tak zwaną gęstość adnotacji (*AnnotationDensity Viewer*), która obrazuje intensywność naniesionych w danym fragmencie wywiadu adnotacji opisujących różne parametry wypowiedzi rozmówcy. Poniżej widoczna jest w formie fali linia czasu wyświetlanego pliku video (*Timeline Viewer*).

Natomiast w dolnej części okna programu znajduje się przeglądarka adnotacji w postaci linii czasu. Dla precyzji działania i lepszego komfortu pracy użytkownika można powiększyć opisywaną linię czasu poprzez wywołanie prawnym przyciskiem myszy menu kontekstowego i wybraniu funkcji powiększania (Zoom).

Adnotacje rozmieszczone są na różnych liniach reprezentujących poszczególne poziomy – *Tiers*). Jednocześnie każdy z poziomów zawiera inny rodzaj adnotacji, na przykład etykietę sekcji, translacji, glosy, etykietę gramatyczną lub fonologiczną, taką jak mrużenie oka lub przesunięcie ciała.

Aktualną pozycję odtwarzania pliku z informacjami o wszystkich adnotacjach i ich nasyceniu (*density*) pokazuje czerwona pionowa linia.

Uwaga: Warto w tym miejscu wspomnieć również o takich funkcjach, jak: *Grid, Text* czy *SubtitleViewers*, które można znaleźć w jednym z menu okna programu. Umożliwiają one różne sposoby przeglądania adnotacji w ramach określonych poziomów: siatkowy (*Grid*), tekstowy (*Text*) oraz w formie napisów (*SubtitleViewers*).



Ilustracja 6.1. Okno programu ELAN

Źródło: opracowanie własne na podstawie programu ELAN

6.4. Rozpoczynanie pracy w programie ELAN

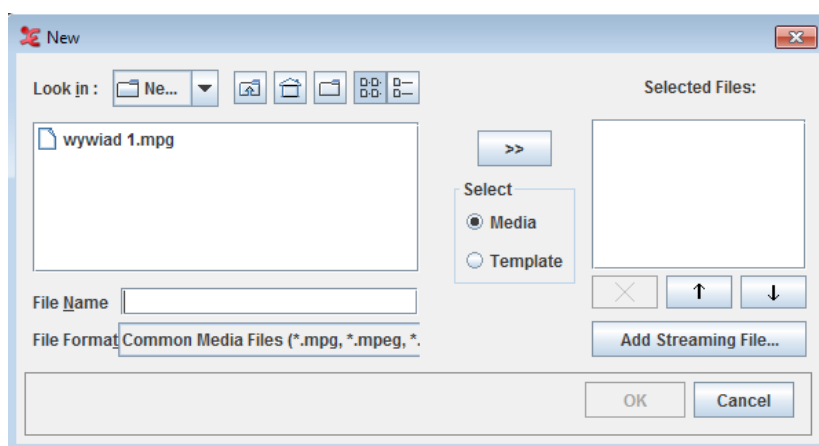
Pracę w programie rozpoczynamy od otwarcia jego okna, co następuje poprzez dwukrotne kliknięcie jego ikony znajdującej na pulpicie komputera (pojawia się automatycznie po instalacji oprogramowania). Po tej czynności należy z menu *File* wybrać jedną z dwóch opcji *Open*, gdy korzystamy z wcześniej utworzonego pliku zapisanego w formacie programu ELAN (*.eaf), lub klikając na *New*, jeśli chcemy zaimportować plik multimedialny (np. *.mpg, *.wav). W tym ostatnim przypadku zostanie wyświetlone kolejne okno dialogowe *New*, w którym należy wykonać następujące czynności:

1. Należy wybrać odpowiedni katalog, w którym znajduje się dany plik multimedialny. Czynność tę wykonujemy poprzez rozwinięcie listy dostępnych katalogów, znajdujących się na komputerze. W tym celu klikamy na pole wyszukiwania *Look in* i przechodzimy do katalogu zawierającego plik multimedialny.

2. Jeśli chcemy używać plików multimedialnych innego typu (np. QuickTime *.mov), wówczas trzeba z menu rozwijanego *File format* wybrać odpowiedni format. Przy czym należy pamiętać, że to, czy dany format będzie obsługiwany, zależy od konfiguracji komputera i zainstalowanego na nim oprogramowania.

3. Następnym krokiem jest dwukrotne kliknięcie na nazwę danego pliku multimedialnego (*.mpg, *.wav itd.), co spowoduje jego pojawianie się w prawym polu okna programu (alternatywnie można wybrać nazwę pliku i kliknąć przycisk >>).

4. Jeśli chcesz używać predefiniowanego zestawu szablonów, wówczas wybieramy przycisk opcji *Temple*, a następnie konkretny szablon (tj. *.etf), który ma zostać użyty.



Ilustracja 6.2. Okno *New* służące do importowania plików multimedialnych w programie ELAN

Źródło: opracowanie własne na podstawie programu ELAN

5. Ostatnią czynnością jest potwierdzenie dokonanych wyborów poprzez kliknięcie na przycisk *OK*, co spowoduje otwarcie nowego dokumentu adnotacji. Aby wyjść z okna dialogowego bez tworzenia nowego pliku, kliknij *Cancel*.

6.5. Podstawowe funkcje programu

Poniżej opisanych zostało kilka podstawowych funkcji umożliwiających swobodne korzystanie z programu ELAN. Pierwsza z nich to nawigowanie plikami multimedialnymi, do którego służy rozbudowany system sterowania w postaci charakterystycznego paska narzędzi (Ilustracja 6.3).

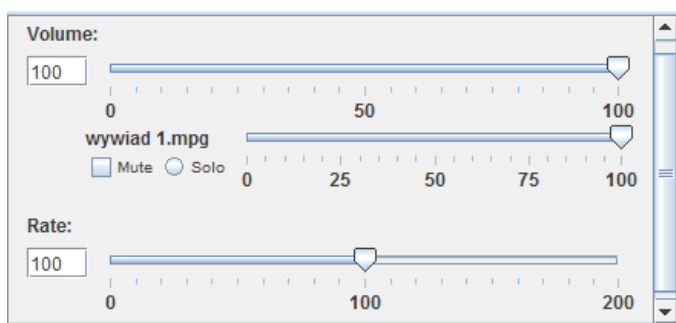


Ilustracja 6.3. Pasek narzędzi służących do sterowania odtwarzanymi plikami multimedialnymi

Źródło: opracowanie własne na podstawie programu ELAN

A zatem w ramach dostępnych opcji nawigowania plikiem multimedialnym można przechodzić do początku lub końca odtwarzanego pliku, przewijać do przodu lub do tyłu o żądaną długość pliku, „przeskakiwać” co sekundę w jedną lub drugą stronę, a także przechodzić o 40 milisekund (domyślnie jedną klatkę pliku) oraz 10 milisekund, co odpowiada jednemu pikselowi na linii czasu (fali) odtwarzanego pliku.

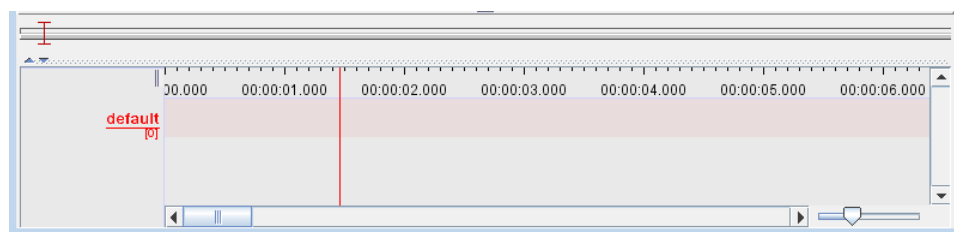
Drugą funkcją jest kontrola prędkości odtwarzania pliku multimedialnego. Program daje możliwość dostosowania prędkości odtwarzania filmu od bardzo powolnej do dwukrotnie szybszej od normalnej. Dotyczy to zarówno ścieżki video, jak i audio.



Ilustracja 6.4. Panel kontroli głośności i prędkości odtwarzanego pliku

Źródło: opracowanie własne na podstawie programu ELAN

Po trzecie każdy plik audio i video może być podzielony na segmenty poprzez wyróżnienie określonych przedziałów na osi czasu. Takie przedziały mogą następnie posłużyć do tworzenia nowych adnotacji lub dokonywania zmian w adnotacjach już istniejących.

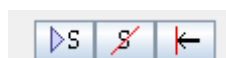


Ilustracja 6.5. Oś czasu z zaznaczonymi przedziałami (segmentami)

Źródło: opracowanie własne na podstawie programu ELAN

W praktyce operacja tworzenia przedziałów polega na zaznaczeniu dowolnego obszaru poprzez przeciągnięcie kursora myszy na osi czasu. Innym sposobem wykonania tej czynności jest przytrzymanie przycisku shift na klawiaturze i kliknięcie myszką na osi czasu w miejscu rozpoczynającym i kończącym obszar, który ma zostać oznaczony, jako oddzielny przedział (segment).

Po czwarte program posiada funkcje, które pozwalają użytkownikowi na swobodne manewrowanie pomiędzy utworzonymi uprzednio segmentami. Służą do tego celu trzy przyciski, znajdujące się pod wyświetlanym plikiem multimedialnym w obszarze paska *Selection*. Pierwszy z nich odpowiedzialny jest za odtwarzanie tylko wybranego segmentu (fragmentu nagrania), a jego połączenie z trybem pętli (*Loop Mode*) spowoduje samoczynne, powtarzające się odtwarzanie nagrania (zamiast przycisku można użyć skrótu klawiaturowego shift + space). Drugi z przycisków powoduje wyczyszczenie zaznaczenia danego fragmentu (ten sam efekt uzyskamy za pomocą klawiszy alt + shift + c lub alt + c. Natomiast trzeci z przycisków daje możliwość przemieszczania się pomiędzy kolejnymi segmentami (fragmentami). Aktywowanie tej funkcji może także nastąpić po użyciu skrótu klawiaturowego ctrl +/-.



Ilustracja 6.6. Przyciski umożliwiające przemieszczanie się pomiędzy utworzonymi segmentami

Źródło: opracowanie własne na podstawie programu ELAN

Piątą funkcją dostępną w programie jest ustalanie trybów pracy. Istnieje kilka różnych trybów pracy, które są zaprojektowane z myślą o konkretnym zadaniu.

Domyślnym trybem jest tryb adnotacji (tryb ogólny), w którym dostępna jest większość funkcji programu. Między innymi dlatego omawiane w tym opracowaniu funkcje są pokazane w trybie adnotacji. Przełączanie pomiędzy trybami można wykonać za pomocą menu *Options* (Opcje).

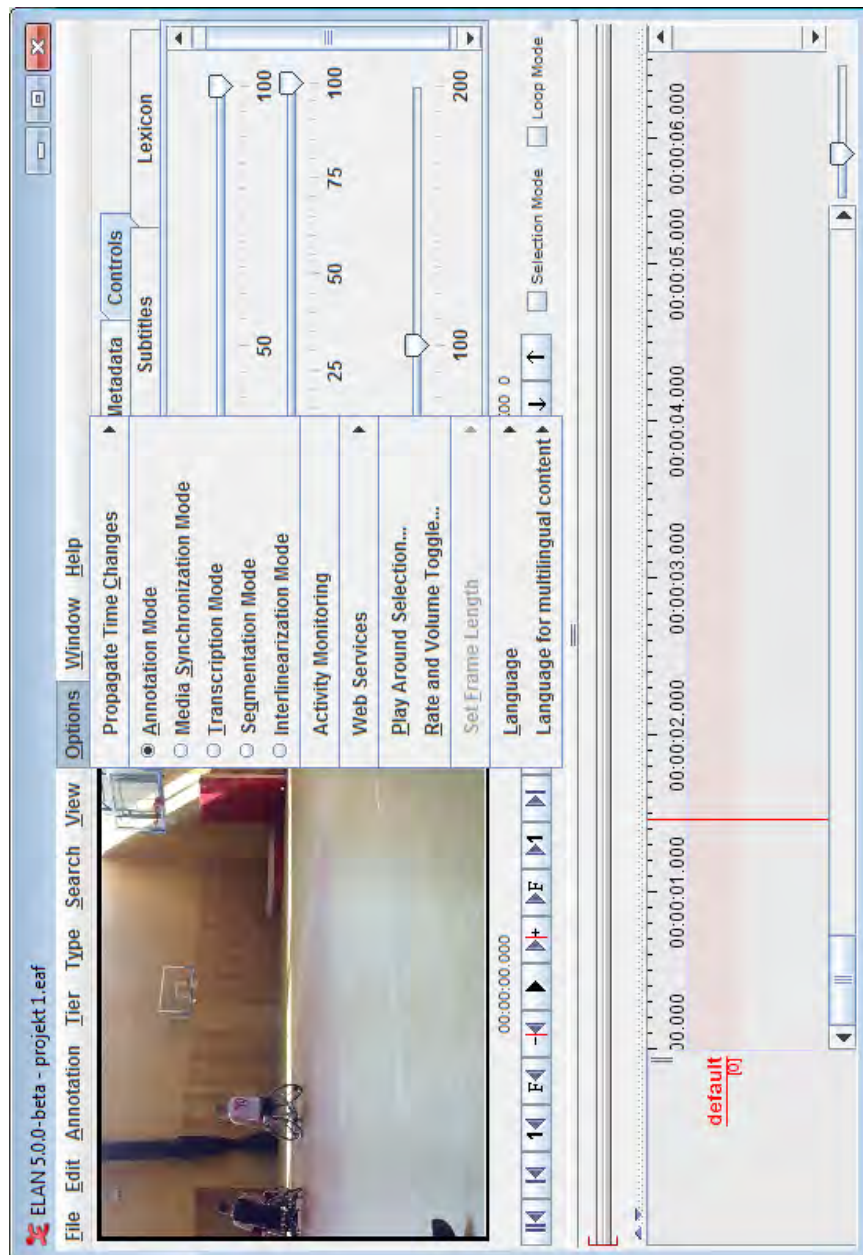
Pozostałe tryby to:

Tryb synchronizacji, który, jak sama nazwa wskazuje, pozwala użytkownikowi na synchronizację plików multimedialnych w sytuacji, gdy poszczególne nagrania nie rozpoczynają się dokładnie w tym samym czasie. Wykorzystując wspomniany tryb pracy, można dokonać przesunięcia w zakresie odtwarzania nagrań, ale bez konsekwencji dla oryginalnych plików multimedialnych. Oznacza to, że synchronizacja wybranych plików będzie widoczna wyłącznie na poziomie programu ELAN.

Tryb transkrypcji jest zoptymalizowany do wpisywania tekstu w istniejące adnotacje. Adnotacje są prezentowane w formie arkusza kalkulacyjnego, w którym nawigacja (z jednej komórki do drugiej) jest całkowicie obsługiwana za pomocą klawiatury. Aktywacja danej komórki (a więc adnotacja), powoduje rozpoczęcie odtwarzania odpowiedniego segmentu (fragmentu) pliku medialnego. Wybór poziomów, które użytkownik chce aktualnie widzieć w tabeli jest oparty na typie warstw, co oznacza, że poziomy tego samego typu są pokazane w tej samej kolumnie.

Tryb segmentacji przeznaczony jest do szybkiego i łatwego tworzenia pustych adnotacji w momencie odtwarzania wybranego nagrania. Tryb ten odznacza się dużą elastycznością w działaniu użytkownika, bowiem nie tylko oznaczenie czasu rozpoczęcia i zakończenia adnotacji odbywa się za pomocą klawiatury, ale też możliwe są zmiany granic adnotacji dokonywane za pomocą prostego przeciągnięcia kursora myszy. Tworzenie i zmiana granic adnotacji w tym trybie nie wymaga uprzedniego dokonywania wyboru czasu (tak jak w trybie opisu).

Tryb interlinearyzacji jest trybem tekstowym, służącym do tworzenia i analizowania adnotacji na jednej lub więcej liniach tekstu. Można to zrobić ręcznie lub za pomocą jednego z tak zwanych analizatorów. Analizatorami są moduły programowe, które przyjmują adnotację jako dane wejściowe i przedstawiają sugestie dotyczące jednej lub więcej adnotacji jako wyjścia. Przykładami typów analizatorów przetwarzania mogą być tokenizacja, analiza morfologiczna i wyszukiwanie glossów. Segmentacja i (zazwyczaj) transkrypcja zdarzeń mowy muszą być dokonane w jednym z pozostałych trybów, zanim będzie można dodać interlinearyzację.



Ilustracja 6.7. Menu *Options* z widocznymi trybami wyświetlania informacji w oknie programu

Źródło: opracowanie własne na podstawie programu ELAN

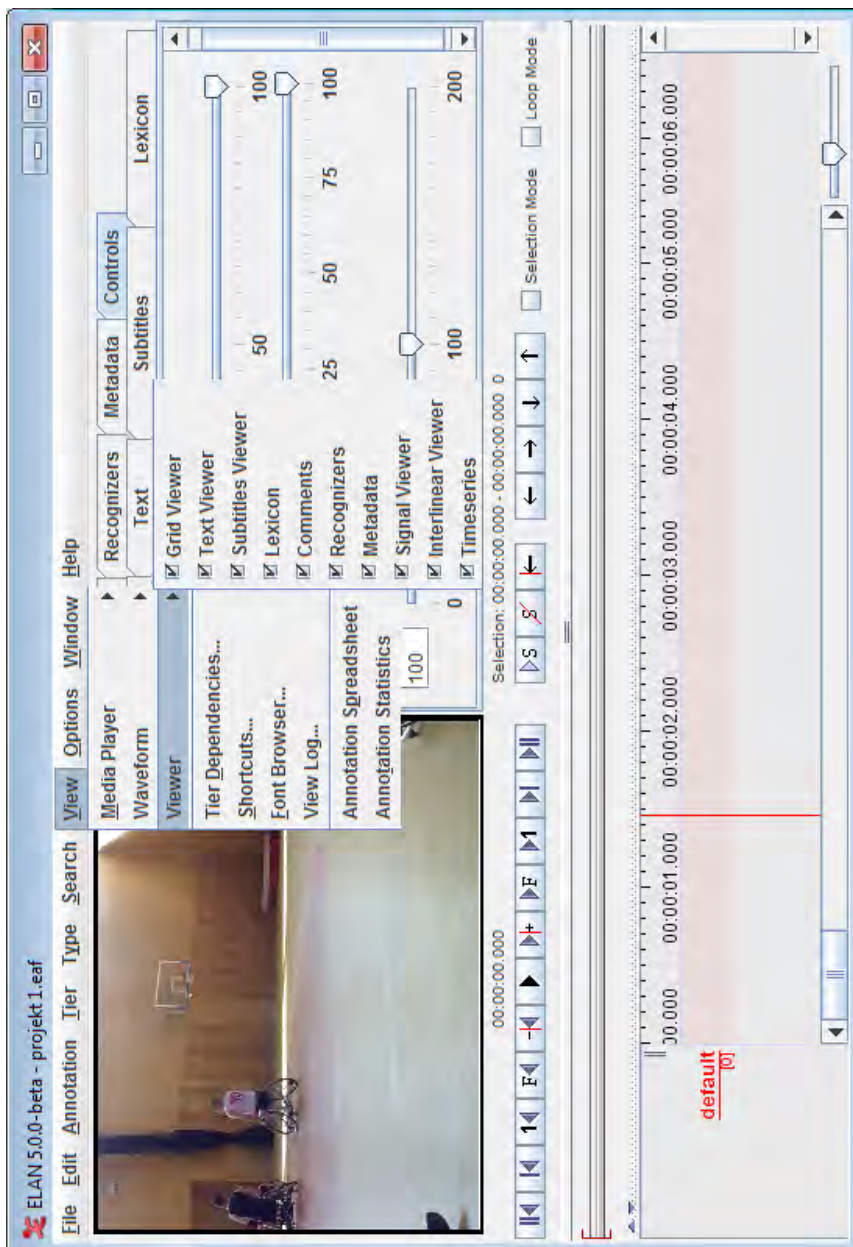
6.6. Funkcje zaawansowane

Program ELAN posiada określoną wewnętrzną architekturę, która wpływa na charakter opisywanego narzędzia. Jako kluczowe elementy wymienia się opisywane już adnotacje, które uszeregowane są w określonego rodzaju warstwy (poziomy). Przy czym program daje możliwość zmiany przyporządkowania danej adnotacji do dowolnego poziomu. Z technicznego punktu widzenia każda taka warstwa (poziom) wyświetlany jest w postaci oddzielnego wiersza adnotacji na linii czasu. Co więcej, można również zmieniać sposób ich prezentacji, wykorzystując w tym celu dostępne sposoby ich wyświetlania w postaci siatki (*Grid*), tekstu (*Text*) oraz nazw (*Subtitles Viewers*).

W obrębie warstwy adnotacje nie nakładają się, ale adnotacje na jednym poziomie mogą się pokrywać z adnotacjami na innych poziomach. Każdy poziom (warstwa) jest również zaklasyfikowany do określonego typu (*Tier Type*). Podstawą ich wyróżnienia są przypisane słowa kluczowe (*Controlled Vocabulary*). W ten sposób można określić, jakie adnotacje mogą być przyporządkowane do danego poziomu (warstwy). Na przykład stopień niepełnosprawności można określić za pomocą słów: „lekki”, „umiarkowany” czy „znaczny”. Byłyby to zatem jedyne możliwe wartości dla adnotacji na tym poziomie.

Jednocześnie, nawet jeśli dwie warstwy lub więcej (poziomów) należą do tego samego typu, to i tak mogą się od siebie różnić na kilka sposobów. Po pierwsze mogą opisać zachowanie różnych uczestników lub różnych artykulatorów tego samego uczestnika. Po drugie mogą być inaczej usytuowane hierarchicznie i przynależać do innych warstw. Po trzecie mogą używać różnych domyślnych języków (angielski, hiszpański, francuski, chiński itp.). I po czwarte mogą używać różnych konwencji lub stylów (np. formalne, a nieformalne głosowanie).

Tak więc na przykład można mieć typ „niepełnosprawności” i mieć dwa różne poziomy, jeden dla dysfunkcji fizycznej, a drugi dla umysłowej. W ten sposób program pozwala na pełną kontrolę nad poziomami, typami i kontrolowanymi słowami używanymi do opisu poszczególnych filmów wideo, zależności między poziomami i sposobami ich prezentacji. Ponadto w tym momencie rozpoczynania pracy przy nowym filmie można wykorzystać jeden z dostępnych szablonów, który zawiera predefiniowany zestaw poziomów, typów poziomów i kontrolowanych słowników. Jednak bardzo często należy dodać nowy poziom lub typ, a także wprowadzić zmianę ich kolejności lub dokonać inne zmiany.



Ilustracja 6.8. Menu View z widocznymi opcjami wyświetlania danych w oknie programu ELAN

Źródło: opracowanie własne na podstawie programu ELAN

6.6.1. Funkcje wyświetlania warstw (poziomów)

ELAN pozwala użytkownikowi na dostosowanie sposobu wyświetlania warstw (poziomów) do jego potrzeb. Program umożliwia bowiem dość swobodne kontrolowanie tego, jakie warstwy i w jakiej kolejności są wyświetlane. Z technicznego punktu widzenia takie zmiany można wykonać na dwa sposoby.

Pierwszy ze sposobów polega na wyborze z menu kontekstowego wyświetlanego za pomocą prawego przycisku myszy, gdy kursor myszy znajduje się na pasku warstw, opcji Sortowania (*Sort Tiers*). W dalszej kolejności musimy określić konkretny sposób wyświetlania warstw.

Sort by Hierarchy – wyświetla poziomy na zasadzie podrzędności/nadrzędności.

Sort by tier Type oraz *Sort by Participant* – grupuje poziomy według typu oraz uczestników.

Unsorted – pozwala na całkowicie dowolną zmianę sposobu wyświetlania warstw.

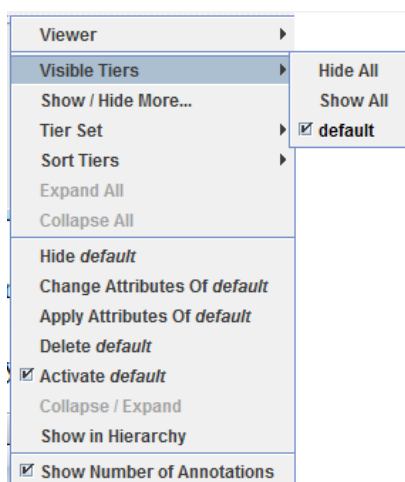
Ilustracja 6.9. Sposoby wyświetlania warstw dostępne w programie ELAN

Źródło: opracowanie własne na podstawie programu ELAN

Drugi sposób zmiany wyświetlania warstw wykorzystuje prostą metodę przeciągania poszczególnych warstw względem siebie – w górę lub w dół na liście. Przy czym w zależności od tego, jaki sposób grupowania poziomów został uprzednio wybrany, będzie to powodowało pewne ograniczenia co do ich zorganizowania. Jeśli na przykład wybrano „Sortuj według uczestnika”, nie będzie można przenieść danego poziomu z dala od innych poziomów, które mają tego samego uczestnika, ale za to można przenieść grupę uczestników do innej grupy w górę lub w dół.

Ponadto dla większej czytelności podglądu warstw w programie wprowadzono możliwość wyświetlania bądź ukrywania całych warstw. Odbывается to poprzez kliknięcie prawym przyciskiem myszy na liście warstw na lewej krawędzi podglądu linii czasowej i wyborze opcji *Visible Tiers*, a następnie jednym kliknięciem na *Hide All* lub *Show All*, co w pierwszym wypadku powoduje ukrycie wszystkich warstw, a w drugim ich wyświetlenie.

Jeśli wybraliśmy sortowanie według hierarchii, to dodatkowo możemy także wyświetlać bądź ukrywać warstwy usytuowane podrzędnie. Wystarczy kliknąć prawym przyciskiem myszy nazwę warstwy nadrzędnej, a następnie wybrać opcję *Zwiń* lub *Rozwiń*.



Ilustracja 6.10. Menu kontekstowe z widocznymi opcjami *Hide All* oraz *Show All*

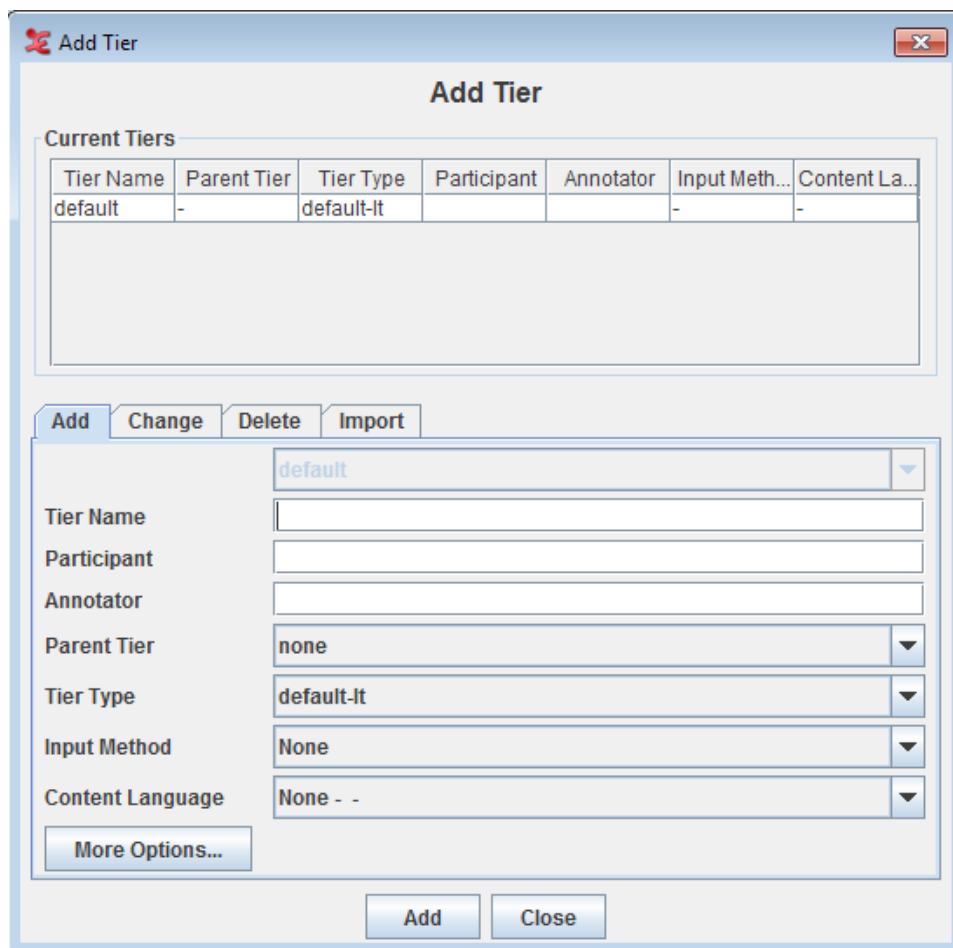
Źródło: opracowanie własne na podstawie programu ELAN

6.6.2. Tworzenie, modyfikowanie i usuwanie warstw

Przed utworzeniem warstw (poziomów) należy w pierwszej kolejności zdecydować o ich organizacji w ramach programu. Będzie to tym bardziej potrzebne, jeśli zdecydujemy się skorzystać z możliwości użytkowania słownika (listy słów kluczowych). Z tego względu ważne będą opcje dodawania nowych warstw adnotacji, ich edycji oraz usuwania. Z technicznego punktu widzenia, aby dodać nową warstwę (poziom), należy wybrać opcję *Add New Tier* z menu *Tier*. Spowoduje to wyświetlenie się nowego okna *Add Tier*, w którym widoczne są wszystkie dotychczasowe warstwy, ale także panel służący ich dodawaniu, edytowaniu i usuwaniu.

W przypadku utworzenia nowego poziomu adnotacji będziemy mogli także wybrać następujące atrybuty:

- nazwę, która powinna być unikalna dla każdego poziomu,
- uczestnika, a więc nazwę lub kod uczestnika, lub podmiotu, którego dotyczy dana kategoria,
- annotatora, czyli nazwę lub kod twórcy adnotacji występujących na tym poziomie,
- jednostkę nadrzędną, o ile chcemy, aby nowo tworzony poziom został przypisany jako podrzędny do innej warstwy,
- typ warstwy, który wskazuje na możliwości i ograniczenia, które mają zastosowanie do niego i zamieszczanych w nim adnotacji,
- domyślny język – w praktyce istotny z perspektywy metod wprowadzania danych, np. za pośrednictwem wirtualnej klawiatury.

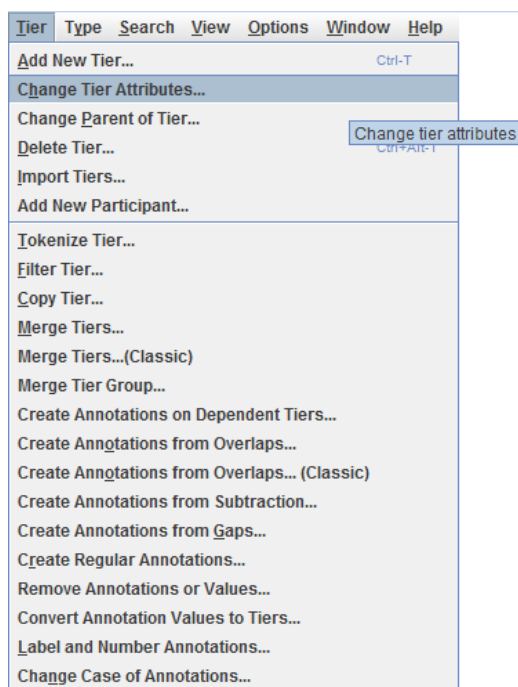


Ilustracja 6.11. Okno *Add Tier* programu ELAN

Źródło: opracowanie własne na podstawie programu ELAN

Jeśli naszym zamiarem będzie modyfikacja istniejących poziomów, wówczas powinniśmy wybierać opcję *Change Tier Attributes* lub kliknąć prawym przyciskiem myszy nazwę warstwy w przeglądarce i wybrać polecenie *Change Attributes* z nazwą danego poziomu. Należy przy tym zaznaczyć, że niektóre zmiany są niedozwolone, ponieważ mogą spowodować usunięcie istniejących adnotacji na danym poziomie.

Aby zaś usunąć warstwę (poziom) trzeba wybrać opcję *Edit*, a następnie *Delete Tier* lub kliknąć prawym przyciskiem myszy nazwę warstwy w podglądzie linii czasu i kliknąć na pozycję *Delete*. Przy czym trzeba pamiętać, że operacja ta spowoduje usunięcie poziomu wraz z adnotacjami, które są na nim związane.



Ilustracja 6.12. Opcja *Change Tier Attributes* dostępna w programie ELAN

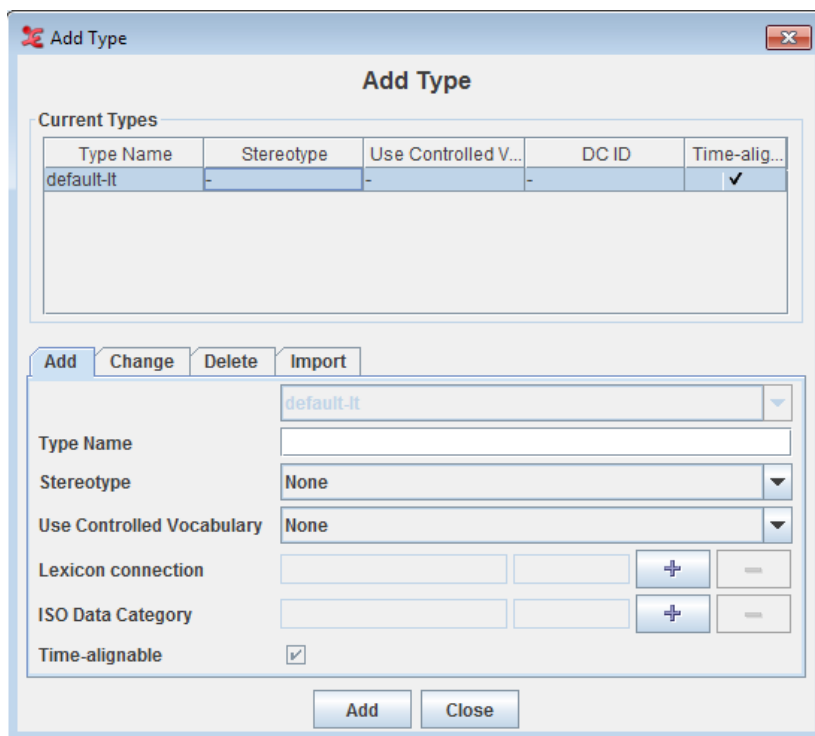
Źródło: opracowanie własne na podstawie programu ELAN

6.6.3. Tworzenie, modyfikowanie i usuwanie typów warstw

Podobnie jak w przypadku samych warstw, tak samo program ELAN pozwala użytkownikowi na działania związane z tworzeniem i edytowaniem ich typów. Dostępnych jest kilka predefiniowanych relacji strukturalnych, które można wykorzystać w procesie tworzenia typów warstw. Jest to o tyle ważne, iż każda warstwa musi być powiązana z określonym typem, a zatem jest przypisana do pewnego rodzaju danych językowych. Z technicznego punktu widzenia, aby dodać nowy typ warstwy, należy wybrać polecenie *Add New Tier Type* z menu the *Type*. Spowoduje to wyświetlenie się okna z bieżącymi typami warstw oraz panelu, gdzie dostępne będą funkcje dodawania, zmiany czy usuwania typów. Dla każdego typu warstw można ustawić odpowiednie atrybuty. Są to:

- nazwa, która podobnie jak w przypadku poszczególnych poziomów powinna być unikalna,
- tak zwane stereotypy, a więc informacje dotyczące jednego z czterech predefiniowanych typów ograniczeń lub ich braku (o czym będzie jeszcze mowa),

- słownictwo kontrolowane, odnoszące się do jednego z dostępnych słowników lub jego braku,
- kategorie danych ISO, czyli identyfikatory kategorii danych w systemie ISOcat,
- „dostrajanie” w czasie, które zależy od wybranego stereotypu i z tego względu nie może być ustawione przez użytkownika.



Ilustracja 6.13. Okno *Tier Type* programu ELAN

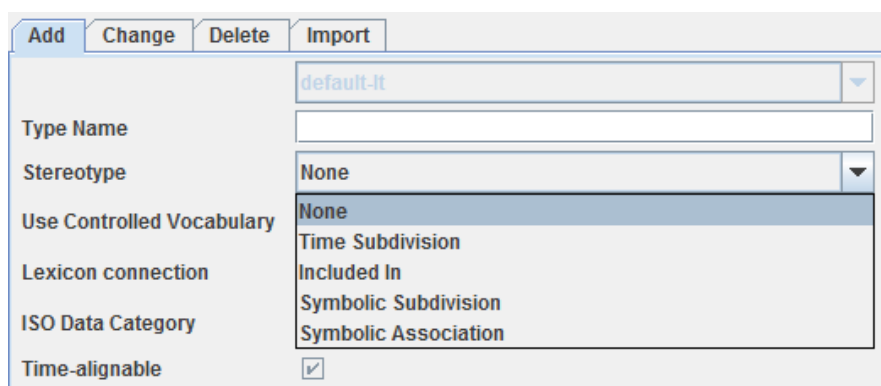
Źródło: opracowanie własne na podstawie programu ELAN

Aby zmodyfikować istniejący typ warstwy, należy kliknąć na polecenie *Edit*, a następnie *Changetier Type*. W otwartym oknie należy wybrać typ, który chce się zmienić, wprowadzić określoną modyfikację i kliknąć na *Change*. Przy czym trzeba pamiętać, że niektóre zmiany są niedozwolone, ponieważ mogą spowodować usunięcie istniejących adnotacji na poszczególnych poziomach. Jeśli zaś chcemy usunąć dany typ warstw, należy wybrać *Edit*, a następnie *Deletetier Type*. W otwartym oknie należy wybrać typ i *Delete*, który chce się usunąć i kliknąć na *Change*. Trzeba przy tym uważać, aby nie usunąć zdefiniowanych poziomów, które należą do danego typu.

6.6.4. Struktura hierarchiczna warstw

Poziomy można ustawiać jako podrzędny lub zależny od innej warstwy. Adnotacje na poziomie zależnym są powiązane z adnotacjami na poziomie nadrzędnym, co oznacza, że granice czasowe adnotacji podrzędnej nie mogą przekraczać granic adnotacji nadrzędnej. Innymi słowy, adnotacje podrzędne zawsze muszą znajdować się w przedziale macierzystej adnotacji. W pierwszej kolejności należy upewnić się, że został utworzony typ warstwy z prawidłowym typem stereotypu.

Aby utworzyć warstwę zależną dla warstw niezależnych, trzeba w oknie *Add Tier* najpierw wybierać poziom nadrzędny dla nowej warstwy, a następnie jeden z dostępnych typów poziomów. Należy pamiętać, że stereotyp decyduje o rodzaju adnotacji, którą można utworzyć na poziomie. Tak długo, jak długo nie zostanie wybrany poziom nadrzędny, wymienione są tylko typy warstw ze stereotypem *None* (brak). Ponadto, podczas ustawiania poziomów obowiązują dodatkowe ograniczenia. Po pierwsze poziom stereotypu *Symbolic Subdivision* lub *Symbolic Association* nie może być nadrzędny wobec typu *Time Subdivision* oraz *Included*. Po drugie po utworzeniu danego poziomu i umieszczeniu w nim adnotacji nie można zmienić typu poziomu na typ z innym stereotypem.



Ilustracja 6.14. Fragment okna *Add Tier* z widocznymi poziomami stereotypów

Źródło: opracowanie własne na podstawie programu ELAN

Poniżej przedstawiono typy stereotypów (ograniczeń) adnotacji i ich najważniejszych cech.

Po pierwsze, to brak, który oznacza, że adnotacja na danym poziomie jest powiązana bezpośrednio z osią czasu, tzn. adnotacja jest wpisana na niezależny poziom, przy czym dwie adnotacje nie mogą się ze sobą pokrywać.

Po drugie, to podział czasu, czyli objaśnienie na poziomie macierzystym można podzielić na mniejsze jednostki, które z kolei mogą być podłączone do prze-

działów czasowych. Przy czym należy pamiętać, że nie mogą występować luki w czasie, co oznacza, że mniejsze jednostki muszą występować bezpośrednio po sobie. Na przykład wypowiedź zapisana na macierzystym szczelbu może być podzielona na słowa, ale każde z nich jest wtedy powiązane z odpowiednim przedziałem czasowym.

Po trzecie, to tak zwany podział symboliczny, który oznacza, że mniejsze jednostki nie mogą być połączone z przedziałem czasowym. Na przykład słowo na rodzimym szczelbu można podzielić na poszczególne morfemy (niezwiązane z czasem).

Po czwarte można wyróżnić ograniczenie „włączenia w” oznaczające, iż wszystkie adnotacje należą do granic poziomu rodzica. Niemniej jednak mogą istnieć luki między adnotacjami. Na przykład, można rozdzielić słowo z ciszą.

Po piąte symboliczne stowarzyszenie, gdzie adnotacja na poziomie rodzica nie może być dzielona na mniejsze składowe i występuje zasada „jeden do jednego” oznaczająca, że na przykład jedno zdanie na poziomie rodzica ma dokładnie jedno tłumaczenie.

6.6.5. Tworzenie i modyfikowanie słowników

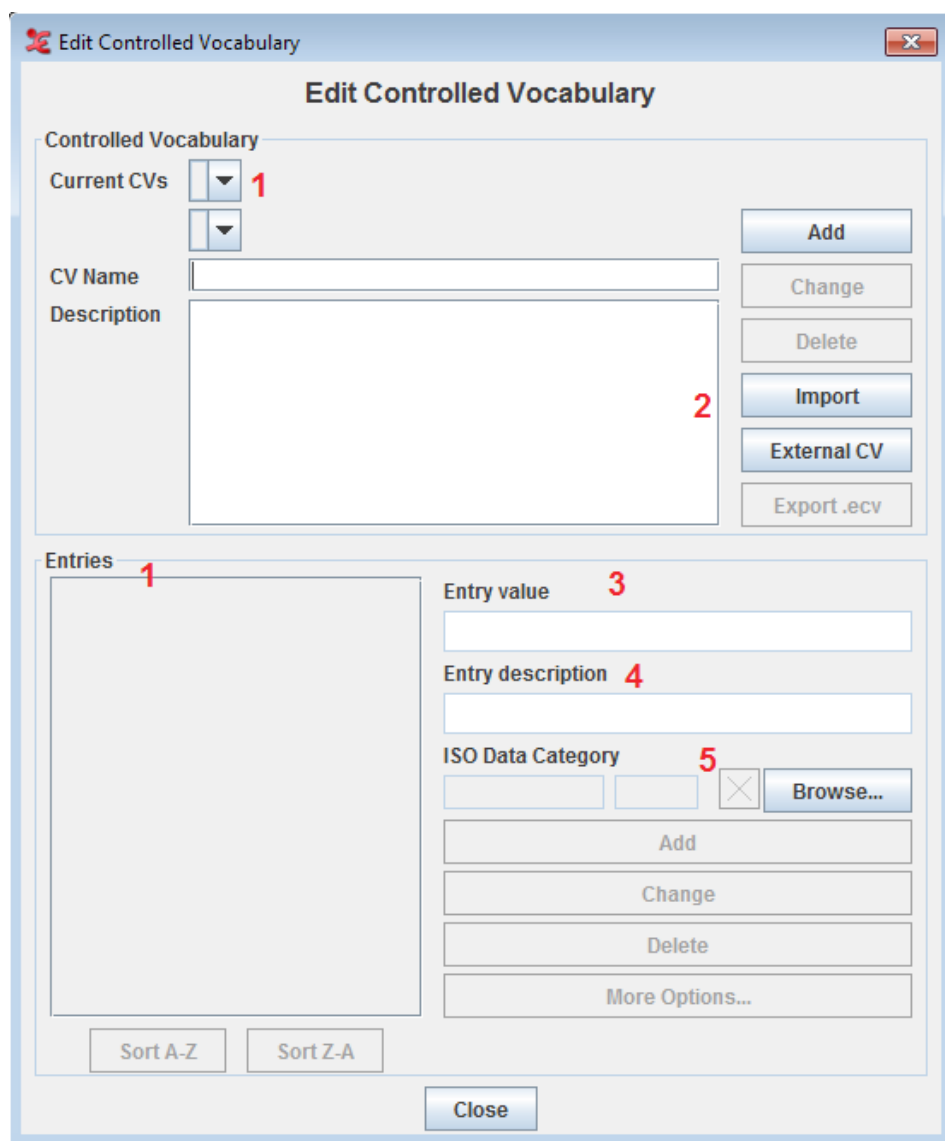
Słownictwo kontrolowane to lista możliwych wartości zapisanych adnotacji, którą użytkownik chce użyć na jednym lub kilku poziomach. Mogą one stanowić część dokumentu adnotacyjnego lub być definiowane w osobnym pliku. Tego rodzaju słownictwo jest sposobem na użycie ograniczonego zestawu etykiet dla adnotacji na jednym poziomie. To zapewnia spójność z pozostałymi adnotacjami i przyspiesza ich dodawanie. Tego rodzaju słownictwo jest szczególnie przydatne do tworzenia adnotacji opisujących cechy fonologiczne, takie jak wzrok, mrugnięcia, kształt ust, przesunięcie ciała itd.

Z technicznego punktu widzenia, aby używać słowników, najpierw należy skonfigurować kontrolowane słownictwo, a następnie utworzyć typ tieru (poziomu), który go używa, i ustawić poziom, który należy do tego typu warstw. Aby zdefiniować nowe słownictwo kontrolowane lub zmienić istniejące w menu *Edit*, należy wybrać polecenie *Controlled Vocabularies*, co spowoduje otwarcie nowego okna.

Okno *Controlled Vocabulary* składa się z dwóch głównych części. Górna część jest przeznaczona do dodawania, importowania, a także zmiany i eksportowania słownictwa. Dolna część okna służy do dodawania i modyfikowania wpisów wybranego słownika. Wpis składa się z wartości i opcjonalnego opisu.

Ponadto opisywane okno zawiera:

1) miejsce, gdzie użytkownik może dokonywać edycji słownictwa, wprowadzać kolejne słowa, ich opisy oraz ustawiać język;



Ilustracja 6.15. Okno *Edit Controlled Vocabularies*

Źródło: opracowanie własne na podstawie programu ELAN

- 2) opcje importu i eksportu słownika, a także możliwość korzystania z zewnętrznych słowników na zasadzie linkowania;
- 3) możliwość wprowadzenia wartości adnotacji, która ma zostać wypełniona;
- 4) opis wprowadzanej wartości (np. części mowy);
- 5) możliwość powiązania z kategorią danych ISO.

Istnieje również możliwość ustawienia jednego lub kilku języków słownika. Pozwala to na przykład na posiadanie jednego słownika zawierającego określone słownictwo zarówno w języku angielskim, jak i polskim. Program pozwala użytkownikowi na wybór spośród prawie 7 700 określonych języków.

W oknie *Edit Controlled Vocabularies* można dokonać wyboru różnych parametrów związanych z danym językiem. Są to:

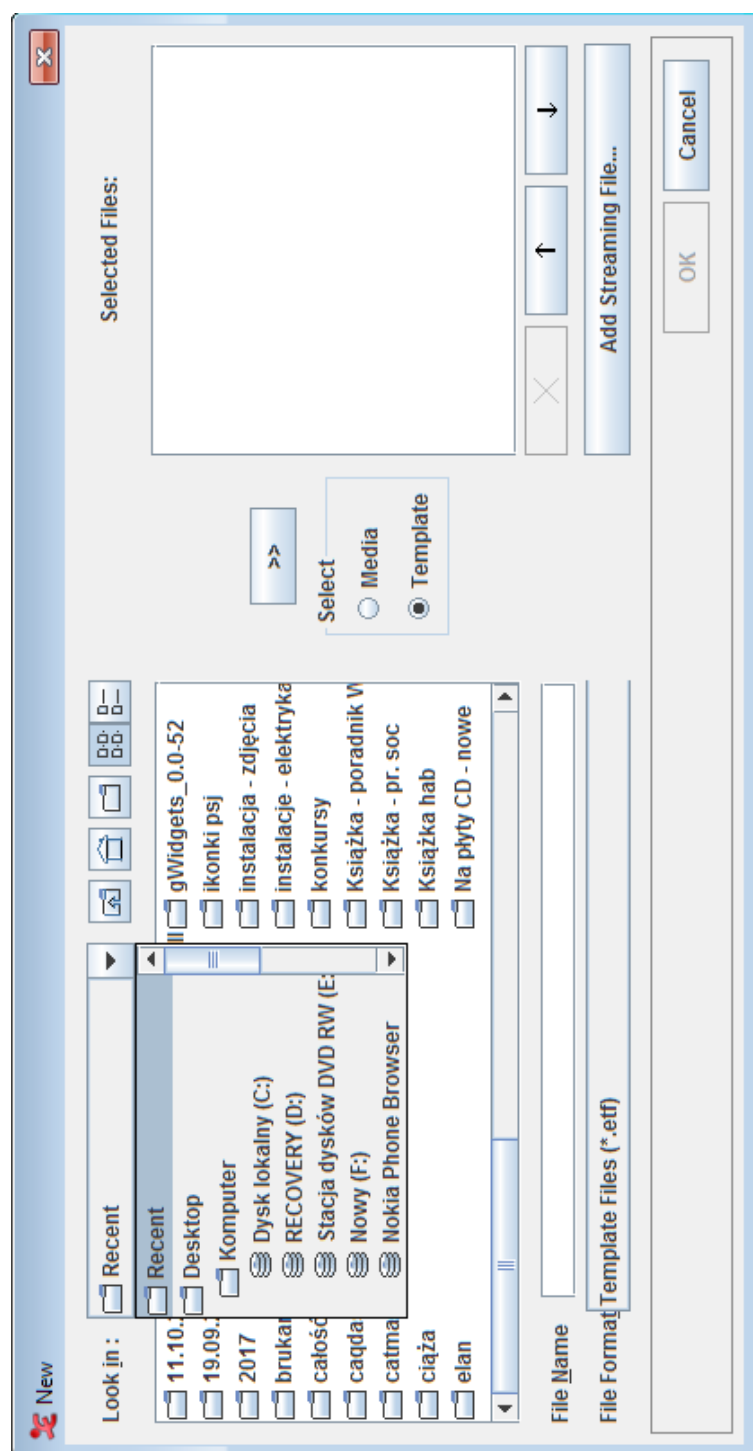
- 1) wybór z listy aktualnych języków, które ma w swojej bazie program ELAN;
- 2) wykonanie wpisu polegającego na powiązaniu (przełożeniu) wybranych języków;
- 3) wprowadzenie wartości opisu;
- 4) przyciski umożliwiające dodawanie, zmianę, usuwanie wpisów w słownikach danego języka.

Warto zaznaczyć, iż w sytuacji, w której poziom jest związany z kontrolowanym słownikiem (za pomocą typu), w chwili tworzenia lub edytowania adnotacji na tym poziomie pojawia się rozwijana lista, pozwalająca użytkownikowi na wybranie jednej z wprowadzonych uprzednio wartości.

6.6.6. Tworzenie i edytowanie szablonów

Program ELAN posiada także wygodną opcję tworzenia i edytowania szablonów. Oznacza to, że gdy użytkownik dokona wyboru poziomów, ich rodzajów, a także słowników może je zapisać i wykorzystać w przyszłości do tworzenia nowych plików adnotacji. Szablon jest dokumentem adnotacyjnym bez linków do plików multimedialnych i bez adnotacji, ale ze zdefiniowanymi poziomami, typami poziomów i słownikiem. Może zostać utworzony z dowolnego dokumentu adnotacyjnego i może służyć jako podstawa do nowych dokumentów adnotacyjnych.

Z technicznego punktu widzenia, aby utworzyć nowy szablon, należy z menu *File* wybrać opcję *Save as Template...* Rozszerzenie pliku, w jakim zapisany jest szablon, to .etf. Aby utworzyć nowy dokument w oparciu o szablon, należy wykonać następujące kroki. W przeglądarce plików klikamy na przycisk *Template* i przechodzimy do folderu zawierającego plik szablonu. Wybieramy jeden z plików zapisanych z rozszerzeniem .etf i klikamy na przycisk >> (copy to right) albo klikamy dwukrotnie na nazwę pliku szablonu. Po wykonaniu tych działań zauważymy, że na prawym panelu znajdzie się teraz plik szablonu i wybrany plik multimedialny dla nowego dokumentu. Nowy dokument jest tworzony z odtwarzaczem multimedialnym i wszystkimi poziomami zdefiniowanymi dla danego szablonu.



Ilustracja 6.16. Okno importowania szablonów w programie ELAN

Źródło: opracowanie własne na podstawie programu ELAN

Należy przy tym pamiętać, że wszelkie zmiany w szablonie nie mają wpływu na tworzone już pliki adnotacji. Z jednej strony chroni to użytkownika przed utratą plików adnotacji, w sytuacji gdy chce się poprawić szablon. Z drugiej strony, jeśli chcemy wprowadzić zmiany w plikach adnotacji, wówczas taką czynność trzeba będzie powtarzać oddzielnie w każdym pliku.

6.6.7. Praca z adnotacjami

„Sercem” programu ELAN jest praca z adnotacjami. W związku z tym w niniejszym punkcie omówione będą kwestie dotyczące dodawania adnotacji do poziomu, a także ich modyfikowania do potrzeb użytkownika.

6.6.7.1. Dodawanie adnotacji

Wprowadzanie adnotacji w programie ELAN może się odbywać na dwa sposoby. Pierwszy z nich opiera się na wpisywaniu tekstu adnotacji w aktywnym polu na obszarze linii czasu, natomiast drugi pozwala na dodawanie adnotacji w oknie edycji adnotacji.

W pierwszym przypadku utworzenie adnotacji zależy od typu warstwy, do której ma zostać dodana adnotacja. Jeśli warstwa, na której ma zostać utworzona adnotacja jest sprzęgnięta z osią czasu, wówczas owe adnotacje są zazwyczaj tworzone z określonym czasem rozpoczęcia i zakończenia danego fragmentu. Natomiast w sytuacji, gdy warstwa nie jest „sparowana” z linią czasu, informacja o długości danego fragmentu nagrania zostanie przeniesiona z opisu adnotacji nadrzędnej.

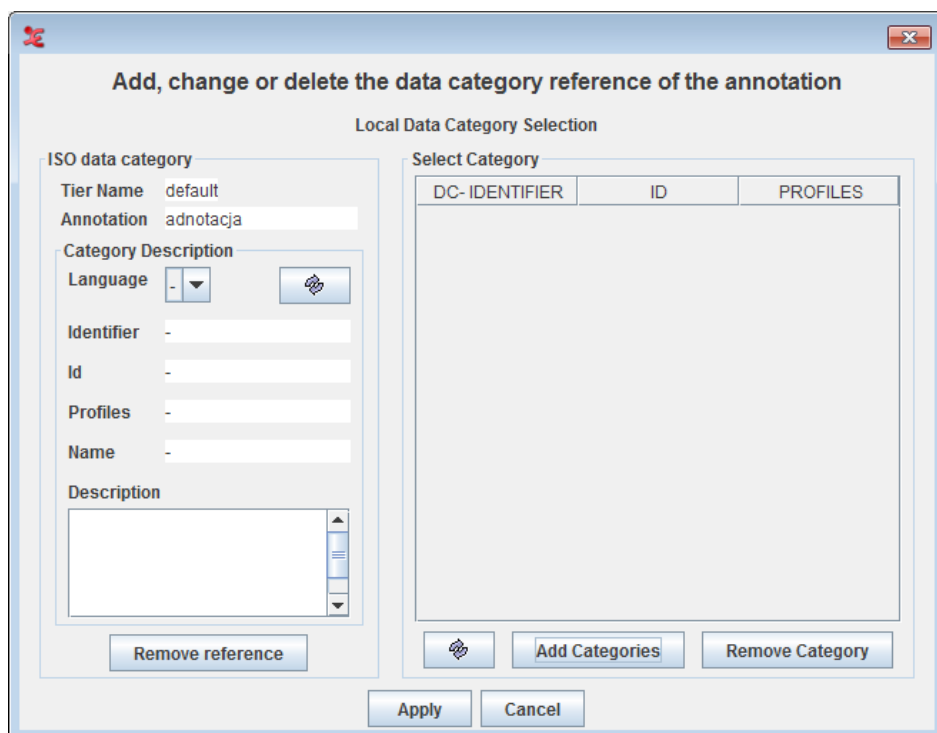
Z technicznego punktu widzenia operacja ta polega na uaktywnieniu danej warstwy pod kątem nanoszenia nowych adnotacji. W tym celu należy wybrać warstwę do pracy jako *Active Tier*. Najłatwiejszym sposobem na to jest dwukrotne kliknięcie nazwy poziomu po lewej stronie podglądu linii czasu. Nazwa aktywnego poziomu zostanie podkreślona na czerwono i pogrubiona. Gdy wybrana została część wideo, zacięniowanie aktywnej warstwy wewnątrz zaznaczenia jest różowe (niebieskie w przypadku plików audio). Aktywną warstwą jest ta, która otrzymuje nowe adnotacje.

Adnotację dostosowaną do czasu można utworzyć na kilka sposobów:

Po pierwsze, w menu *Annotation* należy wybrać pozycję *New Annotation Here*,

Po drugie, należy kliknąć prawym przyciskiem myszy na warstwie i z menu podręcznego wybrać polecenie *New Annotation Here*,

Po trzecie, poprzez dwukrotne kliknięcie w danym obszar, w którym ma być wyrównywany poziom.



Ilustracja 6.17. Okno edycji adnotacji (*Annotation*)

Źródło: opracowanie własne na podstawie programu ELAN

To samo okno pojawia się po dwukrotnym kliknięciu na istniejącą adnotację. Alternatywną metodą wywołania pola edycji jest wykonanie następującego zestawu czynności:

- zaznaczenie czasu, który powinien stać się początkiem adnotacji,
- naciśnięcie klawisza SHIFT + ENTER,
- zaznaczenie czasu, który powinien stać się końcowym czasem adnotacji,
- ponowne naciśnięcie klawisza SHIFT + ENTER.

Wówczas na wybranym poziomie pojawi się pole edycji, gdzie można wprowadzić adnotację, a następnie nacisnąć klawisz ESC, aby zatwierdzić tekst. Jeśli naciśnię się klawisze CTRL + ESC (bez wprowadzania adnotacji), tworzona jest pusta adnotacja. Aby zapisać adnotację, można użyć klawiszy skrótu CTRL + ENTER lub kliknąć prawym przyciskiem myszy w oknie dialogowym Edycja tekstowa i kliknąć przycisk Zatwierdź zmiany w menu rozwijanym.

Drugi sposób dodawania adnotacji polega więc na wprowadzaniu nowego tekstu w specjalnie do tego celu dedykowanym oknie edycji. Okno to różni się od pola edycji tekstowej tym, że obsługuje wprowadzanie dłuższych tekstów.

Pole Edycja adnotacji jest wstępnie skonfigurowane domyślnie. Jeśli chcesz używać innego zestawu znaków, wykonaj następujące czynności:

- a. Kliknij *Select Language*. Pojawi się menu rozwijane zawierające listę zestawów zdefiniowanych znaków;
- b. Kliknij odpowiedni zestaw znaków. Od tej pory znaki są wprowadzane do wybranego zestawu znaków;
- c. Aby powrócić do domyślnego zestawu znaków, powtórz czynności powyżej i wybierz z listy zdefiniowany zestaw.

Po piąte Edytuj adnotację,

Po szóste zapisz adnotację, wykonując jedną z następujących czynności:

- a. Użyj klawiszy skrótu CTRL + ENTER,
- b. W polu Edytuj adnotację kliknij Edytor, a następnie kliknij przycisk Zatwierdź zmiany w menu rozwijanym.

Aby opuścić okno Edycja adnotacji bez zapisywania, wykonaj jedną z następujących czynności:

1. Użyj klawisza skrótu ESC.
2. W oknie Edycja adnotacji kliknij Edytuj, a następnie kliknij polecenie Anuluj zmiany w menu rozwijanym.

Aby powrócić do pola Edycja tekstowa, wykonaj jedną z następujących czynności:

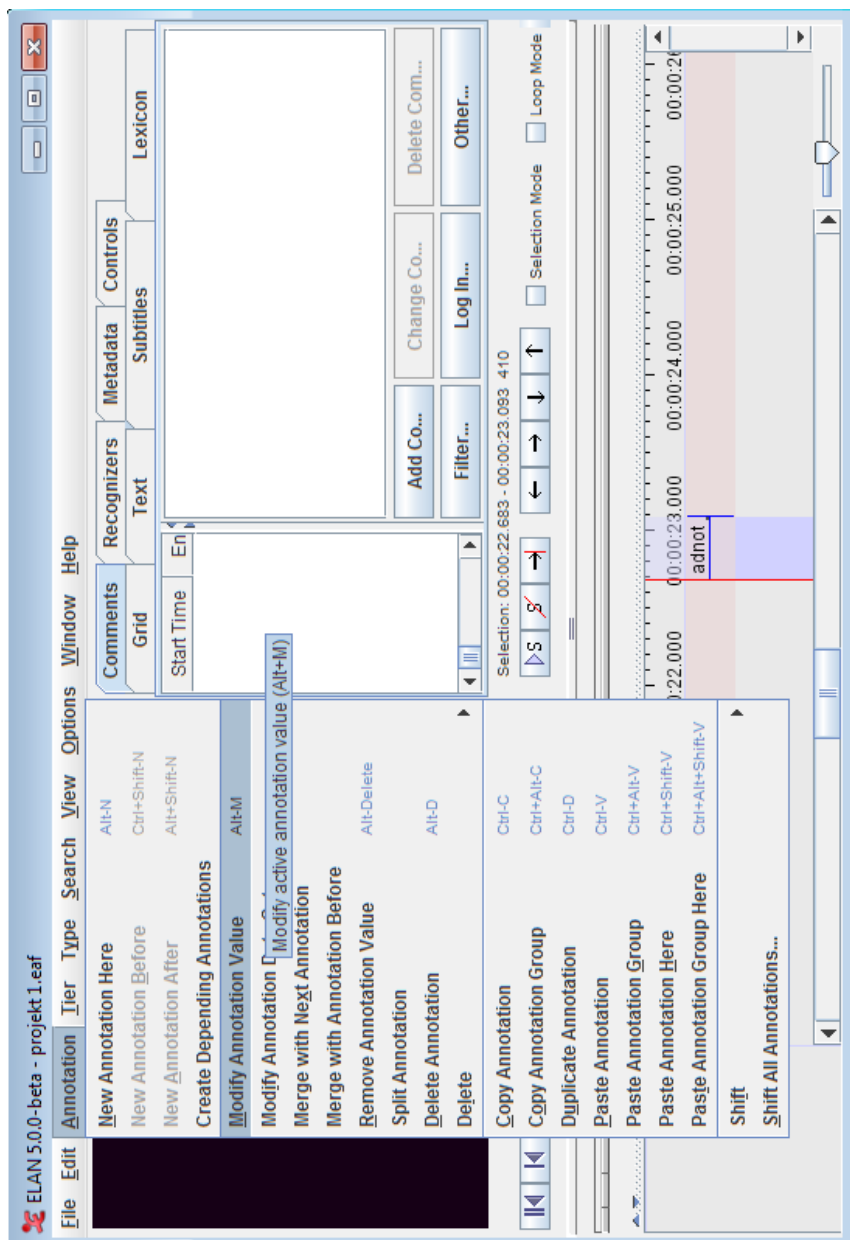
1. Użyj klawiszy skrótu SHIFT + ENTER.
2. W polu Edytuj adnotację kliknij opcję Edytor załączników w menu rozwijanym.

6.6.7.2. Edycja i usuwanie adnotacji

Program ELAN pozwala użytkownikowi na dość swobodne edytowanie adnotacji, co znacznie ułatwia pracę i sprawia, że jest to narzędzie o szerokim spektrum zastosowań. Z technicznego punktu widzenia, aby zmodyfikować tekst istniejącej adnotacji, należy wykonać jedną z następujących czynności.

Na osi czasu lub na podglądzie widoku można kliknąć na wybraną adnotację, co spowoduje pojawianie się ciemnoniebieskiej ramki, co oznacza, że adnotacja jest aktywna. Ewentualnie użytkownik może kliknąć prawym przyciskiem myszy na daną adnotację, co spowoduje pojawienie się rozwijanego menu, w którym należy wybrać opcję *Modify annotation value* (opcję tę wywołać można również, klikając na menu *Annotation*). W obu przypadkach wyświetli się pole edycji (*Edit*).

Alternatywnie ten sam efekt będzie miało dwukrotne kliknięcie adnotacja w *Grid Viewer*, którą chce się zmodyfikować. Wówczas zostanie wyświetlone pole Edycji. Jeśli zaś chcemy usunąć wartość adnotacji, należy wykonać czynności



Ilustracja 6.18. Opcja *Modify annotation value* służąca edycji istniejących adnotacji

Źródło: opracowanie własne na podstawie programu ELAN

tożsame z opisanymi powyżej, a więc na osi czasu lub na podglądzie międzyliniowym można kliknąć na wybraną adnotację, co spowoduje pojawianie się ciemnoniebieskiej ramki, lub klikając prawym przyciskiem, wywołać rozwijane menu, w którym trzeba wybrać opcję *Remove Annotation Value*. Alternatywnie można kliknąć na menu Edycji, a następnie na *Remove Annotation Value*.

6.6.7.3. Dzielenie adnotacji

Adnotacje na niektórych poziomach można podzielić na mniejsze jednostki. W ten sposób można na przykład rozbić zdanie na różne słowa. Aby to uczynić, należy wykonać następujące czynności. Na osi czasu lub na podglądzie międzyliniowym można kliknąć na wybraną adnotację, co spowoduje pojawienie się rozwijanego menu, w którym należy wybrać opcję *New Annotation*. Można również w menu *Annotation* kliknąć opcję *New Annotation*, co spowoduje, że nowa adnotacja pojawi się przed lub po już istniejącej adnotacji. Jeśli klikniemy na opcję *New Annotation*, zostanie ona wstawiona po prawej stronie. Natomiast gdy tę czynność poprzedzimy podświetleniem oryginalnej adnotacji, wówczas zostanie ona podzielona, a nowa adnotacja będzie widoczna po jej lewej stronie. Przy czym należy zauważyć, że opcja ta jest dostępna tylko dla tych warstw, które są przypisane do stereotypów (zdefiniowanych ograniczeń) typu: *Time Subdivision* and *Symbolic Subdivision*. Adnotacja jest zawsze dzielona na dwie jednostki.

6.7. Uwagi końcowe

ELAN to program stworzony z myślą o osobach, które chcą wykonać analizę danych wizualnych z zastosowaniem narzędzi komputerowych. ELAN pozwala tworzyć, edytować, wizualizować i wyszukiwać adnotacje w przypadku danych wideo i audio. Jest to narzędzie przeznaczone specjalnie do analizy języka, w tym także do języka migowego i gestów, ale ma charakter uniwersalny i może być z powodzeniem używane przez wszystkich, którzy pracują z danymi wideo i/lub audio w celach analizy oraz dokumentacji danych. Niewątpliwym atutem tego narzędzia jest dość czytelny interfejs, a także przyjazny dla użytkownika layout. Program wymaga jednak pewnej wprawy, zwłaszcza związanej z prawidłowym rozpoznawaniem pojęć, jakich używa się, operując wspomnianym narzędziem. Niemniej jednak już po kilku próbach można przyzwyczaić się do jego specyfiki i wykonywać różne działania w sposób intuicyjny. Z całą pewnością program ELAN to dobre narzędzie dla tych wszystkich, którzy poszukują alternatywy narzędzi CAQDAS dla drogich programów licencjonowanych, a chcą wykonać w sposób profesjonalny analizę jakościową opartą na danych wizualnych.

Program ELAN został opracowany w Instytucie Psycholingwistyki Maxa Plancka, Nijmegen w Holandii.

Zarówno sam program, jak i dodatkowe informacje o nim można pobrać ze strony:

<http://tla.mpi.nl/tools/tla-tools/elan>

Rozmiar pliku 49,0 MB (wersja 1.2)

Minimalne wymagania to: 256 MB RAM (preferowane 512 MB), procesor 1.5 Ghz. lub szybszy, wersja 32-bit lub 64-bit systemu operacyjnego.

Systemy, pod którymi można uruchomić program: Windows (począwszy od wersji 7), Mac OS X (od wersji v.10.6 lub wyższy) oraz Linux.

ELAN do działania potrzebuje środowiska JAVA.

Zakończenie

CAQDAS jest obecnie reprezentowany przez niezwykle liczną grupę programów różniących się zarówno zaawansowaniem funkcji, jak i ich przeznaczeniem (Gibbs 2011; Richards 2005; por. Prein, Kelle, Bird 1995; Weitzman, Miles 1995). Z tego względu w samej rodzinie oprogramowania CAQDAS występują istotne różnice (Bie-liński, Iwańska, Rosińska-Kordasiewicz 2007: 93; Seale 2008: 233–234; Wilk 2001: 53–56). Obok narzędzi tak rozbudowanych, jak Atlas.ti, NVivo, MaxQDA, QDAMiner i wiele innych, które dają badaczowi możliwość kodowania danych, tworzenia powiązań logicznych i kontekstowych pomiędzy wygenerowanymi kategoriami, czy weryfikowania powstałych hipotez, a w dalszej kolejności także konstruowania teorii (Brosz 2012; Fielding 2007: 463; Kelle 2005: 486; Niedbalski 2014; Niedbalski, Ślęzak 2012: 160), mamy narzędzia prostsze w swej konstrukcji, mniej zaawansowane technicznie i zazwyczaj uboższe w zaimplementowane w nich opcje (Niedbalski 2013a). Nie umniejsza to jednak ich walorów praktycznych i nie powinno stać się powodem negatywnej oceny tego rodzaju narzędzi, zwłaszcza gdy spojrzymy na ceny licencji programów odpłatnych, na które trzeba zazwyczaj przeznaczyć od kilkuset złotych do kilku tysięcy. Mimo iż programy darmowe przeważnie nie są tak rozbudowane, jak dostępne odpłatnie narzędzia, jednak z całą pewnością są wystarczające dla większości, zwłaszcza początkujących badaczy, wyposażając ich w najważniejsze funkcje pomocne w procesie analizy danych jakościowych. Jest to zatem grupa programów, która ma dość szerokie zastosowanie, przy czym dotyczy to w znacznej mierze danych tekstowych i przeważnie ogranicza się do podstawowych, lecz najważniejszych funkcji, takich jak: kodowanie, segregowanie i przeszukiwanie danych. Niemniej jednak, jak starałem się to wykazać w niniejszym opracowaniu, wśród programów bezpłatnych można z powodzeniem odnaleźć także te, które mogą być wykorzystywane do analizy danych wizualnych, w tym materiałów audio i wideo. Niewątpliwie też interesującą i godną uwagi grupę stanowią bezpłatne narzędzia należące do rodziny CAQDAS, a do których można mieć dostęp za pomocą stron internetowych. Są to programy, w wypadku których nie trzeba przechodzić, czasami kłopotliwego, procesu instalacji na własnym komputerze, martwić się o kompatybilność sprzętową czy jego parametry, a jednocześnie wszystkie dane są archiwizowane i przechowywane na zewnętrznych serwerach. Dzięki temu zyskujemy dodatkowe zaplecze dla naszych danych, co stanowi swoisty backup zazwyczaj niezwykle dla nas cennych i pieczołowicie gromadzonych materiałów.

Pisząc kolejną z serii książkę poświęconą CAQDAS, chciałem jeszcze raz zwrócić się przede wszystkim do studentów oraz początkujących badaczy jakościowych, służąc im wskazówkami dotyczącymi tego, jakiego rodzaju oprogramowania mogą użyć, realizując swoje projekty badawcze. Ich ciężka i często długotrwała praca analityczna może stać się bowiem nie tylko bardziej efektywna, ale także po prostu przyjemniejsza. Przy czym nie oznacza to, że będzie prostsza w sensie dosłownym, gdyż usprawnienie samego procesu analizy dotyczy *de facto* strony technicznej i formalnej, nie zaś *stricto* merytorycznej. Z tego względu należy pamiętać, że żaden nawet najbardziej zaawansowany i złożony program komputerowy nie zdejmie z badacza ciężaru pracy analitycznej i interpretacyjnej (zob. Silverman 2007; Travers 2009; Babbie 2006; Konecki 2000).

Na koniec wypada dodać, iż opisane przeze mnie programy w tym jak i wcześniejszych moich opracowaniach w żadnym razie nie wyczerpują bogatej listy tego rodzaju oprogramowania CAQDAS, lecz mogą pomóc rozeznaczyć się w zakresie ich różnorodności i stanowić swego rodzaju przewodnik na drodze do poszukiwań narzędzia najlepszego dla każdego z nas. Jeśli zaś nie mamy pewności, co dokładnie wybrać, powinniśmy po pierwsze zastanowić się nad tym, czego naprawdę potrzebujemy, przetestować różne programy i poważnie zastanowić, jakie funkcje oprogramowania faktycznie wykorzystamy (Niedbalski, Ślęzak 2016; Niedbalski 2013b). Kończąc, życzę, tak jak czynię to za każdym razem, pisać publikacje, wygłaszać prelekcje lub prowadzić warsztaty, owocnych poszukiwań i trafnych wyborów.

Bibliografia

- Babbie Earl (2006), *Badania społeczne w praktyce*, Wydawnictwo Naukowe PWN, Warszawa.
- Bieliński Jacek, Iwańska Katarzyna, Rosińska-Kordasiewicz Anna (2007), *Analiza danych jakościowych przy użyciu programów komputerowych*, „ASK”, 16: 89–114.
- Bringer Joy D., Johnston Lynne H., Brackenridge Celia H. (2006), *Using Computer-Assisted Qualitative Data Analysis Software to Develop a Grounded Theory Project*, „Field Methods”, 18: 245–266.
- Bringer Joy D., Johnston Lynne H., Brackenridge Celia H. (2004), *Maximizing Transparency in a Doctoral Thesis 1: The Complexities of Writing About the Use of QSR*NVIVO Within a Grounded Theory Study*, „Qualitative Research”, 4: 247–265.
- Brosz Maciej (2012), *Zastosowanie pakietu NVivo w analizie materiałów nieustrukturyzowanych* „Przegląd Socjologii Jakościowej” VIII (1), http://www.qualitativesociologyreview.org/PL/Volume18/PSJ_8_1_Brosz.pdf [dostęp: 30.03.2012].
- Charmaz Kathy (2009), *Teoria ugruntowana. Praktyczny przewodnik po analizie jakościowej*, Wydawnictwo Naukowe PWN, Warszawa.
- Charmaz Kathy (2006), *Constructing Grounded Theory. A Practical Guide Through Qualitative Analysis*, Sage Publications, London–New Delhi.
- Coffey Amanda, Atkison Paul (1996), *Making Sense of Qualitative Data: Complementary Research Strategies*, Sage, Thousand Oaks, CA.
- Dahlgren Lars, Emmelin Maria, Winkvist Anna (2007), *Qualitative methodology for international public health*, Epidemiology and Public Health Sciences, Umeå University, Umeå.
- Dohan Daniel, Sanchez-Jankowski Martin (1998), *Using Computers to Analyze Ethnographic Field Data*, „Annual Review of Sociology”, 24: 477–498.
- Fielding Nigel (2007), *Computer Applications in Qualitative Research*, [w:] *Handbook of Ethnography*, P. Atkinson, A. Coffey, S. Delamont, J. Lofland, L. Lofland (eds.), Sage, Los Angeles–London–New Delhi–Singapore.
- Flick Uwe (2010), *Projektowanie badania jakościowego*, Wydawnictwo Naukowe PWN, Warszawa.
- Gibbs Graham (2011), *Analizowanie danych jakościowych*, przeł. Maja Brzozowska-Brywczyńska, Wydawnictwo Naukowe PWN, Warszawa.
- Glaser Barney (1978), *Theoretical Sensitivity*, The Sociology Press, San Francisco.
- Glaser Barney G., Strauss Anselm L. (2009), *Odkrywanie teorii ugruntowanej*, przeł. Marek Gorzko, Zakład Wydawniczy NOMOS, Kraków.
- Glaser Barney G., Strauss Anselm L. (1967), *The discovery of grounded theory. Strategies for qualitative research*, Aldine Publishing Company, Chicago.
- Gorzko Marek (2008), *Procedury i emergencja. O metodologii klasycznych odmian teorii ugruntowanej*, Wydawnictwo Naukowe Uniwersytetu Szczecińskiego, Szczecin.
- Kelle Udo (2005), *Computer-Assisted Qualitative Data Analysis*, [w:] *Qualitative Research Practise*, C. Seale, G. Gobo, J. Gubrium, D. Silverman (eds.), Sage, London–Thousand Oaks–New Delhi.
- Knoblauch Hubert, Flick Uwe, Maeder Christoph (2005), *Qualitative Methods in Europe: The Variety of Social Research* [10 paragraphs], „Forum Qualitative Sozialforschung / Forum: Qualitative Social Research” 6 (3), art. 34, <http://nbn-resolving.de/urn:nbn:de:0114-fqs0503342> [dostęp: 30.03.2012].

- ICT Services and System Development and Division of Epidemiology and Global Health (2013), OpenCode 4.0. Umeå: and Department of Public Health and Clinical Medicine, Umeå University, Sweden, <http://www.phmed.umu.se/enheter/epidemiologi/forskning/open-code/> [dostęp: 10.11.2015].
- Konecki Krzysztof (2012), *Wizualna teoria ugruntowana. Podstawowe zasady i procedury*, „Przegląd Socjologii Jakościowej” VIII (1): 12–45, <http://www.przegladsocjologiijakosciowej.org> [dostęp: 10.11.2015].
- Konecki Krzysztof (2009), *Teaching Visual Grounded Theory*, „Qualitative Sociology Review”, 5 (3), http://www.qualitativesociologyreview.org/ENG/Volume14/QSR_5_3_Konecki.pdf [dostęp: 28.11.2017].
- Konecki Krzysztof (2000), *Studia z metodologii badań jakościowych. Teoria ugruntowana*, Wydawnictwo Naukowe PWN, Warszawa
- Kubinowski Dariusz (2010), *Jakościowe badania pedagogiczne. Filozofia, metodyka, ewaluacja*, Wydawnictwo UMCS, Lublin.
- Lofland John, Snow David A., Anderson Leon, Lofland Lyn H. (2009), *Analiza układów społecznych. Przewodnik metodologiczny po badaniach jakościowych*, przeł. Anna Kordasiewicz, Sylwia Urbańska, Monika Żychlińska, Wydawnictwo Naukowe Scholar, Warszawa.
- Lonkila Marrku (1995), *Grounded theory as an emerging paradigm for computer-assisted qualitative data analysis*: 41–51, [w:] *Computer-Aided Qualitative Data Analysis*, Udo Kelle (ed.), Sage, London.
- Miles Matthew B., Huberman Michael A. (2000), *Analiza danych jakościowych*, Transhumana, Białystok.
- Niedbalski Jakub (2014), *Komputerowe wspomaganie analizy danych jakościowych. Zastosowanie oprogramowania NVivo i Atlas.ti w projektach badawczych opartych na metodologii teorii ugruntowanej*, Wydawnictwo UŁ, Łódź.
- Niedbalski Jakub (2013a), *Odkrywanie CAQDAS. Wybrane bezpłatne programy komputerowo wspomagające analizę danych jakościowych*, Wydawnictwo UŁ, Łódź.
- Niedbalski Jakub (2013b), *CAQDAS – oprogramowanie do komputerowego wspomagania analizy danych jakościowych. Historia, ewolucja i przyszłość*, „Przegląd Socjologiczny”, LXII (1): 153–166.
- Niedbalski Jakub (2012), *OpenCode – narzędzie wspomagające proces przeszukiwania i kodowania danych tekstowych w badaniach jakościowych*, „Przegląd Socjologii Jakościowej” VIII (1), http://www.qualitativesociologyreview.org/PL/Volume18/PSJ_8_1_Niedbalski.pdf [dostęp: 30.03.2012].
- Niedbalski Jakub, Ślęzak Izabela (2016), *Computer Analysis of Qualitative Data in Literature and Research Performed by Polish Sociologists*, „Forum Qualitative Sozialforschung / Forum: Qualitative Social Research (FQS)”, 17 (3): 1–22.
- Niedbalski Jakub, Ślęzak Izabela (2012), *Analiza danych jakościowych przy użyciu programu NVivo a zastosowanie procedur metodologii teorii ugruntowanej*, „Przegląd Socjologii Jakościowej” VIII (1), http://www.qualitativesociologyreview.org/PL/Volume18/PSJ_8_1_Niedbalski_Slezak.pdf [dostęp: 30.03.2012].
- Prein Gerald, Kelle Udo, Bird Katherine (1995), *Computer-Aided Qualitative Data Analysis: Theory, Methods and Practice*, Sage, London.
- Rapley Tim (2010), *Analiza konwersacji, dyskursu i dokumentów*, Wydawnictwo Naukowe PWN, Warszawa.
- Richards Lyn (2005), *Handling Qualitative Data: A Practical Guide*, Sage, London.
- Saillard Elif K. (2011), *Systematic Versus Interpretive Analysis with Two CAQDAS Packages: NVivo and MAXQDA*, „Forum: Qualitative Social Research”, 12 (1), art. 34.
- Seale Clive (2008), *Wykorzystanie komputera w analizie danych jakościowych*, [w:] David Silverman, *Prowadzenie badań jakościowych*, Wydawnictwo Naukowe PWN, Warszawa.

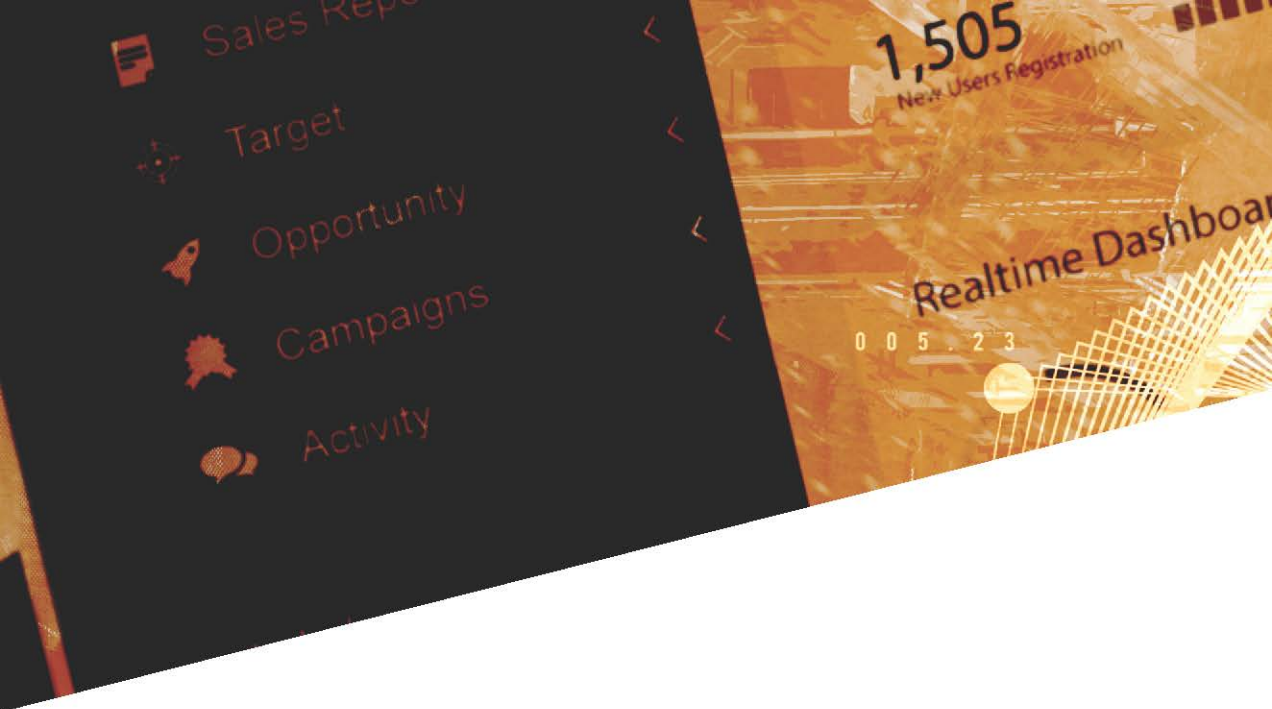
- Silverman David (2008), *Prowadzenie badań jakościowych*, Wydawnictwo Naukowe PWN, Warszawa.
- Silverman David (2007), *Interpretacja danych jakościowych: metody analizy rozmowy, tekstu i interakcji*, Wydawnictwo Naukowe PWN, Warszawa.
- Strauss Anselm L., Corbin Juliet (1990), *Basics of Qualitative Research*, Sage, Thousand Oaks, London–New Delhi.
- Travers Max (2009), *New Methods, Old Problems. A Sceptical View of Innovation*, „Qualitative Research” 9: 161–179.
- Weitzman Eben, Matthew Miles (1995), *Computer Programs for Qualitative Analysis*, Sage, Thousand Oaks, CA.
- Wilk Katarzyna M. (2001), *Komputerowe wspomaganie jakościowej analizy danych*, „ASK”, 10: 49–63.

Introduction to computer analysis of qualitative data. Examples of free CAQDA software. Summary

The goal that I pursue as an author of this book, is to present the possibilities and ways of using selected computer aided qualitative data analysis applications, from the perspective of both a qualitative researcher and user of CAQDA software. This monograph, similarly to the previous one (Niedbalski 2013), is directed first of all to all researchers who are willing to try their hand in their own research projects adopting CAQDA software. In their case, it may turn out to be a good idea to use the free software distributed on a freeware license. Although these applications are usually poorer than their paid versions, as they do not offer so elaborated functions, or they provide less user-friendly environment, in majority they pose a good alternative for still quite expensive – and thus not always widely available – paid applications. And similarly as it happened in case of my first book devoted to CAQDAS (Niedbalski 2013), this time I also decided to present free tools, broadly available via websites, and at the same time fully functional, i.e. without any limitations imposed by their publishers or authors. I reckoned that this is how a user gets an application that can be used in their research without any limits.

Furthermore, the described applications pose a certain selection of possibilities offered by CAQDA software. I intended to present several tools that will not demonstrate an identical set of functions. At the same time, I put great emphasis in this publication not only on the applications allowing to work on a text, but also those that allow to analyze audiovisual materials. I hope that thanks to that the qualitative researchers representing various schools, using various methodologies, will be able to find the tools that suit them.

While selecting particular application, I also followed the utility assets related to specificity of the work environment and their functionality. Therefore, all tools described in the book are characterized with clear structure, architecture and intuitive distribution of particular options. Hence, virtually every person – even with some basic competence in operation of common applications, e.g. word processors – should not have any difficulties in their case.



Jakub Niedbalski, doktor socjologii, adiunkt w Katedrze Socjologii Organizacji i Zarządzania Instytutu Socjologii Uniwersytetu Łódzkiego. Specjalizuje się w komputerowej analizie danych jakościowych, metodach badań jakościowych, zagadnieniach socjologii niepełnosprawności i socjologii sportu. Prowadzi badania dotyczące aktywizacji społecznej i fizycznej osób z niepełnosprawnością. Jest autorem kilkudziesięciu publikacji poświęconych niepełnosprawności, pomocy społecznej oraz metodologii badań jakościowych.

W monografii omówiono programy CAQDA, wspomagające analizę danych jakościowych. Są to narzędzia bezpłatne, powszechnie dostępne, oferujące dość bogate rozwiązania, a niewymagające specjalnych zasobów sprzętowych ani wiedzy z zakresu obsługi standardowych programów komputerowych. Ułatwiają one zarówno pracę z tekstem, jak i analizę materiałów audiowizualnych. Publikacja zawiera też praktyczne wskazówki dotyczące wyboru oprogramowania do realizacji projektów badawczych.

Adresatami książki są nie tylko studenci, doktoranci i badacze zainteresowani analizą danych jakościowych wspartą specjalistycznym oprogramowaniem komputerowym, lecz także wszyscy, którzy w swojej pracy poszukują tego rodzaju narzędzi.

**WYDAWNICTWO
UNIwersytetu
ŁÓDZKIEGO**

 wydawnictwo.uni.lodz.pl

 ksiegarnia@uni.lodz.pl

 (42) 665 58 63

Książka dostępna również
jako e-book

