

User Story

- ❖ The data analysts in our company asked us to prepare waste management datasets and load them to the data warehouse so that they could perform the required analyses
 - ❖ Specifically, they need datasets related to waste management in the Czech Republic at the level of individual regions and the extent of funds used to mitigate the environmental burden due to waste management
 - ❖ Their aim is to assess municipal and corporate waste production and address impacts, and they require that the amount of waste can be calculated per capita or per unit area
 - ❖ As a source of data, we are to use open data on waste from Czech open data portal, data published by the Czech Statistical Office and generally known facts can be extracted from Wikipedia
- ❖ **Data Engineer roles:**
 - ❖ *Extract datasets* from (various) sources
 - ❖ *Transform data* into a uniform form, detect and correct inconsistencies, etc.
 - ❖ *Load* the *data to* a *data warehouse* so that analysts can perform analyses

Prerequisite: Setting up Python (Linux, macOS)

- ❖ *Check* which *version of Python* is installed (if any)
 - ❖ If Python 3 is not installed, download the (latest) version of Python^{#1} and follow the installation
- ❖ Once installed, *create* any *folder for your NDBI046 project* and navigate to it, e.g., ~/Projects/python-ndbi046
- ❖ *Create* your *Python environment*, e.g., ndbi046_env
- ❖ *Activate* you Python environment
- ❖ *Install* the required *packages*
 - ❖ Download the requirements.txt file from the practical class website
 - ❖ You may also *install additional packages*
 - ❖ You may always export the list of installed packages to a file
- ❖ Exit the Python environment after completing the practical class (not before)
 - ❖ You can return to the environment at any time by activating it

^{#1} <https://www.python.org/downloads/>

```
1 % python3 --version
2
3 % mkdir ~/Projects/python-ndbi046
4 % cd ~/Projects/python-ndbi046
5
6 % python3 -m venv ndbi046_env
7
8 % source ndbi046_env/bin/activate
9
10 (ndbi046_env) % pip install -r requirements.txt
11
12 (ndbi046_env) % pip install pandas
13
14 (ndbi046_env) % pip freeze > requirements.txt
15
16 (ndbi046_env) % deactivate
17
18 % cat requirements.txt
```

checking the
installed version

creating
and activating a virtual
environment

installation
of the required
packages

export of
installed packages

list installed
packages

deactivating a
virtual environment

Example 3.2: Transform Region Data (CSV)

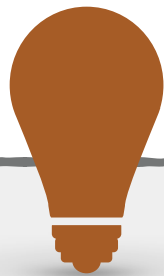
- ❖ Based on the transformation workflow (see Example 3.1), we want to *convert the input dataset* dataset_regions.csv *into* the *dimension* dim_regions
 - ❖ The dimension will contain the attributes region_id (PK), region_name (i.e., region name), and region_area (i.e., region area)
- ❖ Ensure *data consistency* (between datasets)
 - ❖ For example, keep region naming and number formatting consistent across datasets
- ❖ Identify the steps to *transform the dataset* and then *write a Python script* to perform the transformation
 - ❖ Save the transformed result to the file dim_regions.csv

Název kraje (kraj)	Zkratka ČSÚ	Zkratka na RZ	Krajské město	Počet obyvatel	Rozloha (km ²)	Hustota zalidnění	HDP	HDP na obyv.
Hlavní město Praha	PHA	A	Praha	1 357 326	496,21	2 735	1 926,3	1 453,6
Středočeský kraj	STČ	S	Praha	1 439 391	10 928,50	132	775,7	557,6
...	...	...	...	...	...	...	...	...
Zlínský kraj	ZLK	Z	Zlín	580 531	3 963,04	146	304,8	524,9

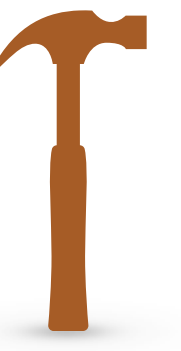


region_id	region_name	region_area
1	Hlavní město Praha	496.21
2	Středočeský kraj	10928.50
...	...	...
14	Zlínský kraj	3963.04

❖ **Tip:** Utilize the Pandas framework for data loading and transformation

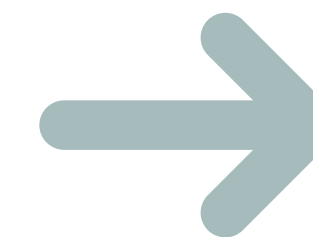


Example 3.2: Transform Region Data (CSV) (Solution)



- ❖ First, we identify the particular steps of the dataset transformation:
 - ❖ Addition of a *surrogate key* (e.g., the index of the row determines the value of the integer key)
 - ❖ *Projection* of columns region_id, Název kraje (kraj), and Rozloha (km²)
 - ❖ *Renaming* columns Název kraje (kraj) → region_name a Rozloha (km²) → region_area
 - ❖ Modification of the *number formatting* in the region_area column from Czech to English format

Název kraje (kraj)	Zkratka ČSÚ	Zkratka na RZ	Krajské město	Počet obyvatel	Rozloha (km ²)	Hustota zalidnění	HDP	HDP na obyv.
Hlavní město Praha	PHA	A	Praha	1 357 326	496,21	2 735	1 926,3	1 453,6
Středočeský kraj	STČ	S	Praha	1 439 391	10 928,50	132	775,7	557,6
...	...	...	...	...	...	...	...	...
Zlínský kraj	ZLK	Z	Zlín	580 531	3 963,04	146	304,8	524,9

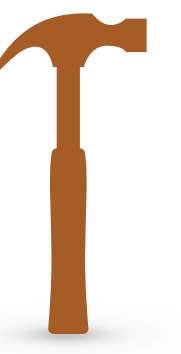


an integer primary key

projection and column name change

region_id	region_name	region_area
1	Hlavní město Praha	496.21
2	Středočeský kraj	10928.50
...	...	...
14	Zlínský kraj	3963.04

adjusted number formatting



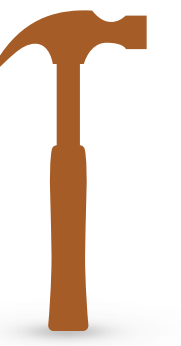
Example 3.2: Transform Region Data (CSV) (Solution)

```
1 import logging
2 import sys
3 import pandas as pd
4
5 logging.basicConfig(level = logging.INFO, format="%levelname)s: %(message)s")
6
7 def load_csv_to_dataframe(file_path: str) -> pd.DataFrame:
8     pass
9
10 def add_surrogate_key(dataframe: pd.DataFrame, ident_name: str) -> pd.DataFrame:
11     pass
12
13 def project_columns(dataframe: pd.DataFrame, columns: list) -> pd.DataFrame:
14     pass
15
16 def rename_columns(dataframe: pd.DataFrame, column_names: dict) -> pd.DataFrame:
17     pass
18
19 def format_numbers(dataframe: pd.DataFrame, column: str) -> pd.DataFrame:
20     pass
21
22 def save_dataframe_to_csv(dataframe: pd.DataFrame, file_path: str) -> None:
23     pass
```

Import the
necessary libraries

Logging settings

Decompose the
task into individual
subtasks with
appropriate
interface

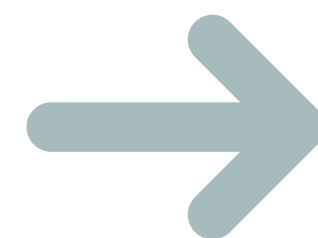


Example 3.2: Transform Region Data (CSV) (Solution)

```
7 def load_csv_to_dataframe(file_path: str) -> pd.DataFrame:
8     try:
9         dataframe = pd.read_csv(file_path)
10        return dataframe
11    except Exception as e:
12        logging.error(f"Error loading CSV file: {e}")
13        raise
14
15 def add_surrogate_key(dataframe: pd.DataFrame, ident_name: str) -> pd.DataFrame:
16     try:
17         dataframe[ident_name] = range(1, len(dataframe) + 1)
18         return dataframe
19     except Exception as e:
20         logging.error(f"Error adding surrogate key: {e}")
21         raise
```

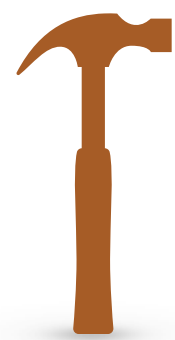
loading a csv file
using pandas

adding an integer
surrogate key, where the row
index determines its value



region_id	Název kraje (kraj)	Zkratka ČSÚ	Zkratka na RZ	Krajské město	Počet obyvatel	Rozloha (km ²)	Hustota zalidnění	HDP	HDP na obyv.
1	Hlavní město Praha	PHA	A	Praha	1 357 326	496,21	2 735	1 926,3	1 453,6
2	Středočeský kraj	STČ	S	Praha	1 439 391	10 928,50	132	775,7	557,6
...	...	...	...	...	...	...	...	...	...
14	Zlínský kraj	ZLK	Z	Zlín	580 531	3 963,04	146	304,8	524,9

Example 3.2: Transform Region Data (CSV) (Solution)

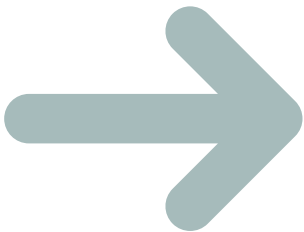


```
23 def project_columns(dataframe: pd.DataFrame, columns: list) -> pd.DataFrame:
24     try:
25         return dataframe[columns]
26     except Exception as e:
27         logging.error(f"Error projecting columns: {e}")
28         raise
29
30 def rename_columns(dataframe: pd.DataFrame, column_names: dict) -> pd.DataFrame:
31     try:
32         return dataframe.rename(columns=column_names, copy=False)
33     except Exception as e:
34         logging.error(f"Error renaming columns: {e}")
35         raise
```

the operator [] performs the projection of the specified columns

rename the columns

region_id	Název kraje (kraj)	Zkratka ČSÚ	Zkratka na RZ	Krajské město	Počet obyvatel	Rozloha (km ²)	Hustota zalidnění	HDP	HDP na obyv.
1	Hlavní město Praha	PHA	A	Praha	1 357 326	496,21	2 735	1 926,3	1 453,6
2	Středočeský kraj	STČ	S	Praha	1 439 391	10 928,50	132	775,7	557,6
...	...	...	...	...	...	...	...	...	...
14	Zlínský kraj	ZLK	Z	Zlín	580 531	3 963,04	146	304,8	524,9



region_id	region_name	region_area
1	Hlavní město Praha	496,21
2	Středočeský kraj	10 928,50
...	...	...
14	Zlínský kraj	3 963,04

