

PYTHON DATA CLEANING SCRIPTS

1

Cleaning of OWID Database *compact.csv*

```
1. # ETL Cleaning Script for OWID Datasets (compact.csv and
   vaccinations_manufacturer.csv)
2. # Author: Miguel Atarés
3. # This script performs data cleaning and preprocessing on OWID datasets using
   Python and pandas.

4. import pandas as pd

5. df = pd.read_csv(r'C:\Users\miguel\Desktop\COVID-19
   Project\Databases\compact.csv', encoding = "ISO-8859-1", delimiter=',')
6. df.head(5)
7. df.shape
8. df.describe()
9. df.info()
10. # Lo primero es eliminar las columnas que ya he planeado que no quiero
11. df2 = df.drop(['code', 'new_cases', 'new_cases_per_million',
   'new_deaths_smoothed',
   'new_deaths_smoothed_per_million', 'hosp_patients', 'weekly_hosp_admissions',
   'weekly_hosp_admissions', 'icu_patients_per_million', 'weekly_icu_admissions',
   'total_tests', 'new_tests', 'new_tests_per_thousand', 'total_vaccinations',
   'people_vaccinated', 'people_fully_vaccinated', 'new_vaccinations',
   'total_vaccinations_per_hundred', 'people_vaccinated_per_hundred',
   'total_boosters_per_hundred', 'new_vaccinations_smoothed_per_million',
   'new_people_vaccinated_smoothed', 'new_people_vaccinated_smoothed_per_hundred'],
   axis = 1)
12. df2.info() #Todos los valores numéricos están como floats, no hay errores de
   columnas en string, solo tipo date que no es datetime
13. df2.isnull().mean().sort_values(ascending=False) #Me da % de valores nulos. A
   más alto, más valores nulos

14. # Vemos que hay 2 columnas con 100% de nan y es un valor que se puede buscar y
   poner manualmente, pero otros tienen % alto de nans no se porque aun
15. df2.shape
16. # Vamos a eliminar duplicados:
17. df2.drop_duplicates(inplace=True)
18. df2.shape #Vemos que no hay duplicados
19. lista_paises = df2["country"].unique()
20. lista_paises
21. len(lista_paises)
22. # Muchos "paises" no son ni paises, y otros me entorpecen el analisis pues se
   detectan muchos nan al ser países poco desarrollados y no reportar datos
```

```

23. # Vamos a eliminar los paises que no nos interesan, ya que no tienen datos
    suficientes, quedandome con los paises mas desarrollados y con más datos

24. latam_top10 = [
25. 'Argentina', 'Brazil', 'Chile', 'Colombia', 'Mexico', 'Peru',
26. 'Uruguay', 'Costa Rica', 'Ecuador', 'Panama'
27. ]

28. africa_top10 = [
29. 'Nigeria', 'South Africa', 'Egypt', 'Morocco', 'Algeria',
30. 'Kenya', 'Ethiopia', 'Ghana', 'Ivory Coast', 'Tunisia'
31. ]

32. asia_top10 = [
33. 'China', 'India', 'Japan', 'South Korea', 'Indonesia', 'Iran',
34. 'Saudi Arabia', 'Israel', 'Vietnam', 'Thailand'
35. ]

36. norteamerica = [
37. 'Canada', 'United States'
38. ]

39. europa = [
40. 'United Kingdom', 'Germany', 'France', 'Italy', 'Spain', 'Netherlands',
41. 'Belgium', 'Sweden', 'Switzerland', 'Poland', 'Portugal', 'Austria',
42. 'Ireland', 'Norway', 'Finland', 'Denmark', 'Czechia', 'Hungary',
43. 'Greece', 'Slovakia', 'Slovenia', 'Romania', 'Bulgaria', 'Croatia',
44. 'Estonia', 'Lithuania', 'Latvia', 'Luxembourg', 'Malta'
45. ]

46. oceania = [
47. 'Australia', 'New Zealand', 'Fiji', 'Papua New Guinea'
48. ]

49. # Unimos todos los países a conservar en una sola lista
50. paises_validos = set(latam_top10 + africa_top10 + asia_top10 + norteamerica +
    europa + oceania)

51. df3 = df2[df2['country'].isin(paises_validos)].copy() #Con esto creamos un nuevo
    df con los paises interesantes
52. df3.shape #Hemos reducido muchísimo el dataset, de 506841 a 128002 filas.
53. df3['country'].unique() #Comprobamos que todo este ok.
54. len(df3['country'].unique()) #Ahora solo tenemos 64 paises, y columnas reducidas
    claro.
55. # Ahora vamos a ver cuantos nan hay en total y comparamos
56. df3.isnull().mean().sort_values(ascending=False) #Sigue habiendo muchas columnas
    con alto % de NaNs.

```

```

57. # Vamos a ver que países no reportan datos en ninguna columna para ver si los
    podemos descartar
58. # Lista de columnas con alto porcentaje de nan (excluyendo las totalmente
    vacías)
59. cols_a_comprobar = [
60.     "excess_mortality",
61.     "excess_mortality_cumulative",
62.     "excess_mortality_cumulative_absolute",
63.     "excess_mortality_cumulative_per_million",
64.     "weekly_icu_admissions_per_million",
65.     "weekly_hosp_admissions_per_million",
66.     "total_boosters",
67.     "hosp_patients_per_million",
68.     "icu_patients",
69.     "people_fully_vaccinated_per_hundred",
70.     "handwashing_facilities",
71.     "total_tests_per_thousand",
72.     "tests_per_case",
73.     "new_tests_smoothed_per_thousand",
74.     "new_tests_smoothed",
75.     "positive_rate",
76.     "new_vaccinations_smoothed",
77.     "reproduction_rate",
78.     "stringency_index"
79. ]

80. # Iteramos por cada columna y mostramos los países que tienen todo nan en esa
    columna
81. for col in cols_a_comprobar:
82.     paises_nunca_reportan = df3.groupby("country")[col].apply(lambda x:
        x.isna().all())
83.     paises_nunca_reportan =
        paises_nunca_reportan[paises_nunca_reportan].index.tolist()
84.     print(f"\n🚫 Países que NUNCA reportan datos en: {col}")
85.     print(paises_nunca_reportan if paises_nunca_reportan else "Todos los países
        tienen al menos algún dato.")

86. # Con esto ya sabemos países con muchos huecos en los datos. sin embargo,
    decido no eliminar más países. lo tendré en cuenta para la visualización en
    Power BI
87. df3["date"].sort_values(ascending=False) #Revisamos fechas y veo que hay fechas
    que ni si quiera han ocurrido aun... Vamos a eliminarlas
88. # Primero, convertimos la fecha a datetime64[ns] de pandas:
89. # Convertir la columna 'date' a tipo datetime
90. df3["date"] = pd.to_datetime(df3["date"], errors='coerce')

91. # Eliminar filas con fecha posterior al 1 de junio de 2025
92. df3 = df3[df3["date"] <= "2025-06-01"]

```



```

121. "Finland": 82.0, "France": 82.5, "Germany": 80.6, "Ghana": 63.8, "Greece": 80.1,
122. "Hungary": 74.5, "India": 67.2, "Indonesia": 67.6, "Iran": 73.9, "Ireland":
    82.0,
123. "Israel": 82.3, "Italy": 82.9, "Japan": 84.8, "Kenya": 61.4, "Latvia": 73.6,
124. "Lithuania": 73.7, "Luxembourg": 82.6, "Malta": 83.8, "Mexico": 70.2, "Morocco":
    74.0,
125. "Netherlands": 81.7, "New Zealand": 82.5, "Nigeria": 52.7, "Norway": 83.2,
    "Panama": 76.2,
126. "Papua New Guinea": 65.4, "Peru": 72.4, "Poland": 76.5, "Portugal": 81.0,
    "Romania": 74.2,
127. "Saudi Arabia": 76.9, "Slovakia": 74.9, "Slovenia": 80.7, "South Africa": 62.3,
128. "South Korea": 83.7, "Spain": 83.0, "Sweden": 83.0, "Switzerland": 84.0,
    "Thailand": 78.7,
129. "Tunisia": 73.8, "United Kingdom": 80.7, "United States": 77.2, "Uruguay": 75.4,
130. "Vietnam": 73.6
131. }

132. manos = {
133. "Algeria": 0.85, "Argentina": 0.95, "Australia": 0.99, "Austria": 0.99,
    "Belgium": 0.99,
134. "Brazil": 0.86, "Bulgaria": 0.95, "Canada": 0.99, "Chile": 0.70, "China": 0.90,
135. "Colombia": 0.70, "Costa Rica": 0.86, "Croatia": 0.90, "Czechia": 0.99,
    "Denmark": 0.94,
136. "Ecuador": 0.87, "Egypt": 0.90, "Estonia": 0.99, "Ethiopia": 0.08, "Fiji": 0.95,
137. "Finland": 0.99, "France": 0.99, "Germany": 0.99, "Ghana": 0.42, "Greece": 0.99,
138. "Hungary": 0.90, "India": 0.38, "Indonesia": 0.79, "Iran": 0.95, "Ireland":
    0.99,
139. "Israel": 0.99, "Italy": 0.99, "Japan": 0.99, "Kenya": 0.38, "Latvia": 0.99,
140. "Lithuania": 0.99, "Luxembourg": 0.99, "Malta": 0.92, "Mexico": 0.86, "Morocco":
    0.84,
141. "Netherlands": 0.99, "New Zealand": 0.99, "Nigeria": 0.37, "Norway": 0.94,
    "Panama": 0.87,
142. "Papua New Guinea": 0.13, "Peru": 0.70, "Poland": 0.99, "Portugal": 0.90,
    "Romania": 0.95,
143. "Saudi Arabia": 0.97, "Slovakia": 0.99, "Slovenia": 0.90, "South Africa": 0.44,
144. "South Korea": 0.95, "Spain": 1.00, "Sweden": 0.99, "Switzerland": 0.99,
    "Thailand": 0.39,
145. "Tunisia": 0.84, "United Kingdom": 1.00, "United States": 1.00, "Uruguay": 0.98,
146. "Vietnam": 0.40
147. }

148. # Aplicar los valores manuales al dataframe. con .map(diccionario)
149. df3["human_development_index"] = df3["country"].map(idh)
150. df3["life_expectancy"] = df3["country"].map(vida)
151. df3["handwashing_facilities"] = df3["country"].map(manos)
152. # Revisamos a ver si es correcto
153. df3.groupby("country")[["human_development_index", "life_expectancy",
    "handwashing_facilities"]].mean().reset_index()

154. df3.isnull().mean().sort_values(ascending=False)

```

```

155. # Vamos a ver porque hay tantos nan en algunas variables como mortalidad. Parece
    que columnas como las de excess_mortality reportan solo 1 vez al mes, por ello
    hay mucho % de NaNs
156. resultado = df3.groupby('country')['excess_mortality'].mean().reset_index()

157. # Mostrar todas las filas una por una
158. for index, row in resultado.iterrows():
159. print(row)

160. # Me da media casi todos los valores, eso significa que en este caso, no es
    porque ciertos paises nunca reporten, si no que algunos paises han debido
    reportar en algun momento, pero pocas veces
161. # Se me ha olvidado quitar valores anteriores a la pandemia. Datos empiezan
    01/01/2020 y se declaro en 11/03/2020, por lo que establecemos eliminar desde el
    01/03/2020
162. df3 = df3[df3["date"] >= "2020-03-01"]
163. df3.shape #Hemos reducido de 126656 a 122816, perfecto.
164. # Por último, reviso columnas categóricas
165. resultado =
    df3.groupby("country")["continent"].nunique().sort_values(ascending=False)

166. for country, num_continentes in resultado.items():
167. print(f"{country}: {num_continentes}")
168. # ¿Cada pais tiene correctamente asociado su continente? Lo revisamos con este
    código:
169. agrupado = df3.groupby("country")["continent"].unique()

170. for country, continents in agrupado.items():
171. print(f"{country}: {list(continents)}")

172. # Perfecto, todas dan un solo valor de continente y he revisado que esté bien,
    por lo que todo ok
173. df3.to_csv("df3_limpio_para_powerbi.csv", index=False)
174. # Vale, me faltaba hacer un describe() para revisar valores que tiene cada
    columna ya que muchos pueden ser ilógicos
175. pd.set_option('display.max_columns', None)

176. df3.describe()
177. # 1. tests_per_case:
178. # Problema: la mediana es ≈17, pero el máximo es 789.603, lo cual sugiere que
    por cada positivo hubo casi 790.000 test... absurdo
179. # Umbral lógico: más de 500 tests por caso se considera probablemente erróneo
180. # 2. reproduction_rate:
181. # Problema: mínimo negativo (-0.056) y máximo 4.64. el valor negativo no tiene
    sentido epidemiológicamente
182. # Umbral lógico: no debería ser menor que 0

183. # 3. otras columnas (no limpiaremos más)

```

```

184. # Por ejemplo, new_cases_smoothed de 5 millones en un dia, ratios de excess
    mortality muy elevados o negativos, etc. hay muchos valores de estos
185. # Que considerados sospechosos. pero haciendo un poco de research, sabemos que
    hay días en los que el covid pegó muy fuerte, por lo que puede
186. # Ser que haya habido registros así en ciertas fechas, al final una pandemia
    puede escalar exponencialmente en ciertos momentos
187. # Vamos a definir condiciones lógicas para seguir limpiando los datos y ver si
    hay errores en los valores extremos
188. errores = {
189. "tests_per_case": df3["tests_per_case"] > 500,
190. "reproduction_rate": df3["reproduction_rate"] < 0,
191. }

192. # Ver cuántas filas están afectadas por cada condición
193. for col, condition in errores.items():
194. print(f"{col}: {condition.sum()} filas afectadas")

195. # Ver ejemplos de los valores problemáticos
196. for col, condition in errores.items():
197. print(f"\nEjemplos de errores en {col}:")
198. display(df3[condition][["country", "date", col]].head(5))

199. # Crear copia limpia y reemplazar los valores extremos por nan
200. df4 = df3.copy()
201. for col, condition in errores.items():
202. df4.loc[condition, col] = pd.NA

203. # Verificación final de limpieza aplicada
204. print("Recuento tras limpieza:")
205. print("Reproduction rate negativos:", df4["reproduction_rate"].lt(0).sum())
206. print("Tests per case mayores a 500:", df4["tests_per_case"].gt(500).sum())

207. df4.describe() #Comprobamos que este todo ok conforme lo que queriamos.
208. # Con esto ya hemos acabado el analisis exploratorio y podemos empezar en power
    bi con este dataset. Repasemos lo hecho o no hecho aquí:
209. # ☒ 2.1) limpiar los nan
210. # ☒ 2.2) formato datetime
211. # ☐ 2.3) modificar decimales --> en pbi
212. # ☒ 2.4) revisar duplicados
213. # ☒ 2.5) revisar columnas categóricas
214. # ☒ 2.6) outliers (iqr, boxplots) --> outliers y medidas extrañas comprobados
    con describe()
215. # ☒ 2.7) distribución de datos (histogramas, frecuencias) --> en pbi
216. # ☒ 2.8) filtrar fechas previas a la pandemia
217. # ☒ 2.9) correlaciones --> realizaremos en power bi
218. # ☒ 2.10) eliminar columnas innecesarias
219. # ☒ 2.11) revisar valores de desviación como mínimos y máximos de columnas

220. df4.to_csv("df4_limpio_para_powerbi.csv", index=False)

```

```

221. # Ahora empezamos la limpieza de otra base de datos: vaccinations manufacturer
222. # Solo hay que poner los mismos paises, controlar las mismas fechas y revisar
    columna de tipos de vacuna
223. df_vacc = pd.read_csv(r'C:\Users\miguel\Desktop\COVID-19
    Project\Databases\vaccinations_manufacturer.csv', encoding = "ISO-8859-1",
    delimiter=',')
224. df_vacc.shape
225. df_vacc.info()
226. df_vacc2 = df_vacc[df_vacc['country'].isin(paises_validos)].copy()
227. df_vacc2.shape
228. df_vacc2['country'].unique()
229. # Esto significa que faltan algunos países en el dataset, simplemente porque no
    habría datos de muchos de ellos
230. # Además el número de filas no ha bajado mucho, por lo que originalmente en el
    dataset había pocos países, muchos ya incluidos en nuestra lista

231. # Esta es la lista de paises que nos faltan en este dataset y están en el
    primero:
232. # Algeria, australia, brazil, china, colombia, costa rica, egypt, ethiopia,
    fiji, ghana, greece, india, indonesia, iran, israel, kenya, mexico,
233. # Morocco, new zealand, nigeria, panama, papua new guinea, saudi arabia,
    thailand, tunisia, united kingdom, vietnam

234. # Ahora transformamos la fecha y excluimos las que no nos interesan:
235. # Convertir la columna 'date' a tipo datetime
236. df_vacc2["date"] = pd.to_datetime(df_vacc2["date"], errors='coerce')

237. # Filtrar filas solo entre el 1 de marzo de 2020 y el 1 de junio de 2025 (ambos
    inclusive)
238. df_vacc2 = df_vacc2[(df_vacc2["date"] >= "2020-03-01") & (df_vacc2["date"] <=
    "2025-06-01")]

239. df_vacc2.info()
240. df_vacc2.shape
241. df_vacc2['vaccine'].unique() #Vemos como está categorizado la columna vaccine
242. # Eliminamos los duplicados
243. df_vacc2.drop_duplicates(inplace=True)
244. df_vacc2.shape #No hay duplicados
245. # Me interesa crear una nueva columna para poder asociar cada nombre de vaccine
    con el tipo de tecnología que usa la propia vacuna
246. # Diccionario de mapeo vacuna → tecnología
247. vaccine_technology = {
248. "Pfizer/BioNTech": "ARNm",
249. "Moderna": "ARNm",
250. "Oxford/AstraZeneca": "Vector viral",
251. "Johnson&Johnson": "Vector viral",
252. "Sputnik V": "Vector viral",

```

```

253. "Sputnik Light": "Vector viral",
254. "CanSino": "Vector viral",
255. "Sinovac": "Virus inactivado",
256. "Sinopharm/Beijing": "Virus inactivado",
257. "Covaxin": "Virus inactivado",
258. "Valneva": "Virus inactivado",
259. "Novavax": "Proteína recombinante",
260. "Sanofi/GSK": "Proteína recombinante",
261. "SKYCovione": "Proteína recombinante",
262. "Medicago": "Otras"
263. }

264. # Crear nueva columna 'technology' con el tipo de tecnología correspondiente
265. df_vacc2["technology"] = df_vacc2["vaccine"].map(vaccine_technology)

266. # Revisamos que se haya añadido correctamente:
267. print(df_vacc2[["vaccine", "technology"]].drop_duplicates().sort_values("vaccine"
    ))
268. # Vemos que en total hay 5 categorías nuevas
269. df_vacc2['technology'].unique()
270. # Vamos a revisar los NaN ahora:
271. df_vacc2.isnull().mean().sort_values(ascending=False)
272. # Perfecto, no hay ningún nan. ya podemos utilizar la database
273. df_vacc2.to_csv("df_vacc2_limpio_para_powerbi.csv", index=False)

```

2

OWID Database vaccinations_*manufacturer.csv*:

```

1. # Ahora empezamos la limpieza de otra base de datos: vaccinations manufacturer
2. # Solo hay que poner los mismos países, controlar las mismas fechas y revisar
   column de tipos de vacuna
3. df_vacc = pd.read_csv(r'C:\Users\migue\Desktop\COVID-19
   Project\Databases\vaccinations_manufacturer.csv', encoding = "ISO-8859-1",
   delimiter=',')
4. df_vacc.shape
5. df_vacc.info()
6. df_vacc2 = df_vacc[df_vacc['country'].isin(países_validos)].copy()
7. df_vacc2.shape
8. df_vacc2['country'].unique()
9. # Esto significa que faltan algunos países en el dataset, simplemente porque no
   habría datos de muchos de ellos
10. # Además el número de filas no ha bajado mucho, por lo que originalmente en el
    dataset había pocos países, muchos ya incluidos en nuestra lista

11. # Esta es la lista de países que nos faltan en este dataset y están en el
    primero:
12. # Algeria, australia, brazil, china, colombia, costa rica, egypt, ethiopia,
    fiji, ghana, greece, india, indonesia, iran, israel, kenya, mexico,
13. # Morocco, new zealand, nigeria, panama, papua new guinea, saudi arabia,
    thailand, tunisia, united kingdom, vietnam

```

```

14. # Ahora transformamos la fecha y excluimos las que no nos interesan:
15. # Convertir la columna 'date' a tipo datetime
16. df_vacc2["date"] = pd.to_datetime(df_vacc2["date"], errors='coerce')

17. # Filtrar filas solo entre el 1 de marzo de 2020 y el 1 de junio de 2025 (ambos
    inclusive)
18. df_vacc2 = df_vacc2[(df_vacc2["date"] >= "2020-03-01") & (df_vacc2["date"] <=
    "2025-06-01")]

19. df_vacc2.info()
20. df_vacc2.shape
21. df_vacc2['vaccine'].unique() #Vemos como está categorizado la columna vaccine
22. # Eliminamos los duplicados
23. df_vacc2.drop_duplicates(inplace=True)
24. df_vacc2.shape #No hay duplicados
25. # Me interesa crear una nueva columna para poder asociar cada nombre de vaccine
    con el tipo de tecnología que usa la propia vacuna
26. # Diccionario de mapeo vacuna → tecnología
27. vaccine_technology = {
28.     "Pfizer/BioNTech": "ARNm",
29.     "Moderna": "ARNm",
30.     "Oxford/AstraZeneca": "Vector viral",
31.     "Johnson&Johnson": "Vector viral",
32.     "Sputnik V": "Vector viral",
33.     "Sputnik Light": "Vector viral",
34.     "CanSino": "Vector viral",
35.     "Sinovac": "Virus inactivado",
36.     "Sinopharm/Beijing": "Virus inactivado",
37.     "Covaxin": "Virus inactivado",
38.     "Valneva": "Virus inactivado",
39.     "Novavax": "Proteína recombinante",
40.     "Sanofi/GSK": "Proteína recombinante",
41.     "SKYCovione": "Proteína recombinante",
42.     "Medicago": "Otras"
43. }

44. # Crear nueva columna 'technology' con el tipo de tecnología correspondiente
45. df_vacc2["technology"] = df_vacc2["vaccine"].map(vaccine_technology)

46. # Revisamos que se haya añadido correctamente:
47. print(df_vacc2[["vaccine",
    "technology"]].drop_duplicates().sort_values("vaccine"))
48. # Vemos que en total hay 5 categorías nuevas
49. df_vacc2['technology'].unique()
50. # Vamos a revisar los NaN ahora:
51. df_vacc2.isnull().mean().sort_values(ascending=False)
52. # Perfecto, no hay ningún nan. ya podemos utilizar la database
53. df_vacc2.to_csv("df_vacc2_limpio_para_powerbi.csv", index=False)

```

3

Cleaning of VAERS Database *data.csv*

```
1. # Author: Miguel Atarés
2. # coding: utf-8
3. # Análisis de la Base de Datos: DATA2021
4. # </h1>

5. import pandas as pd

6. df_VAERS_2021_Data = pd.read_csv(r'C:\Users\migue\Desktop\COVID-19
Project\Databases\2021VAERSData\2021VAERSDATA.csv', encoding = "ISO-8859-1",
delimiter=',')

7. #Nos da un error que indica que algunas columnas son inconsistentes en sus datos
(mezclan tipos de datos)
8. #Para solucionarlo, le vamos a decir que ignore este mensaje y cargue todo
normalmente. Ya revisaremos luego los tipos de datos.

9. df_VAERS_2021_Data = pd.read_csv(
10. r'C:\Users\migue\Desktop\COVID-19
Project\Databases\2021VAERSData\2021VAERSDATA.csv',
11. encoding="ISO-8859-1",
12. delimiter=',',
13. low_memory=False
14. )

15. df_VAERS_2021_Data.shape

16. df_VAERS_2021_Data.info()

17. #Lo primero es eliminar columnas que ya analizamos previamente que no nos
interesan

18. # Lista definitiva de columnas que quieres conservar
19. columnas_a_conservar = [
20. "VAERS_ID", "RECVDATE", "VAX_DATE", "ONSET_DATE", "NUMDAYS", "DATEDIED",
21. "STATE", "AGE_YRS", "SEX",
22. "DIED", "L_THREAT", "ER_VISIT", "HOSPITAL", "HOSPDAYS",
23. "DISABLE", "RECOVD", "BIRTH_DEFECT"
24. ]

25. # Filtrar el DataFrame para conservar solo las columnas deseadas
26. df2_VAERS_2021_Data = df_VAERS_2021_Data[columnas_a_conservar].copy()

27. df2_VAERS_2021_Data.shape

28. #Pasamos de (768712, 36) a (768712, 17)

29. df2_VAERS_2021_Data.drop_duplicates(inplace=True) #Eliminamos duplicados
```

```

30. df2_VAERS_2021_Data.shape #Había algun duplicado como podemos observar

31. df2_VAERS_2021_Data.info()

32. df2_VAERS_2021_Data.head()

33. #Vamos a pasar a datetime las fechas, que son 4 columnas en total:
34. fechas = ["RECVDATE", "VAX_DATE", "ONSET_DATE", "DATEDIED"]
35. for col in fechas:
36.     df2_VAERS_2021_Data[col] = pd.to_datetime(df2_VAERS_2021_Data[col],
        errors="coerce")

37. #Las demás columnas están en el formato correcto de datos

38. categoricas = [
39.     "SEX", "DIED", "L_THREAT", "ER_VISIT", "HOSPITAL",
40.     "DISABLE", "RECOVD", "BIRTH_DEFECT", "STATE"
41. ]

42. for col in categoricas:
43.     print(f"\n🔍 Valores únicos en la columna '{col}':")
44.     print(df2_VAERS_2021_Data[col].unique())

45. # Hay inconsistencia en dos columnas. La primera RECOVD con 'U'. Suele ser
    'unknown', pero lo vamos a pasar a NaN para que no de problemas
46. # En STATE hay algunos estados en minustuculas (Ca, Tx, Co) y 4 códigos no
    documentados en VAERS (XB, XL, XV, QM)
47. # El resto de estados de STATE son los 50 americanos, Washington y otros
    territorios como puerto rico, guam, islas virgenes, etc

48. #Primero sustituimos U por NaN:
49. df2_VAERS_2021_Data["RECOVD"] = df2_VAERS_2021_Data["RECOVD"].replace("U",
    pd.NA)

50. #Ahora solucionamos el problema con la columna STATE:

51. #1. Corrección de valores sin todo mayusculas
52. df2_VAERS_2021_Data["STATE"] = df2_VAERS_2021_Data["STATE"].str.upper()

53. # 2. Conservamos solo los estados válidos
54. estados_validos = [
55.     'AL', 'AK', 'AZ', 'AR', 'CA', 'CO', 'CT', 'DE', 'FL', 'GA', 'HI', 'ID', 'IL', 'IN', 'IA', 'KS',
        'KY', 'LA', 'ME', 'MD', 'MA', 'MI',
56.     'MN', 'MS', 'MO', 'MT', 'NE', 'NV', 'NH', 'NJ', 'NM', 'NY', 'NC', 'ND', 'OH', 'OK', 'OR', 'PA',
        'RI', 'SC', 'SD', 'TN', 'TX', 'UT',
57.     'VT', 'VA', 'WA', 'WV', 'WI', 'WY', 'DC', 'PR', 'VI', 'GU', 'AS', 'MP', 'MH', 'FM'
58. ]

```

```

59. df2_VAERS_2021_Data.loc[~df2_VAERS_2021_Data["STATE"].isin(estados_validos),
    "STATE"] = pd.NA

60. #¿Por qué aparece Puerto Rico y otros similares en datasets como VAERS?
61. #   Puerto Rico (PR), Guam (GU), Islas Vírgenes (VI), Samoa Americana (AS),
    Islas Marianas del Norte (MP), entre otros,
62. #   son territorios no incorporados de Estados Unidos. No son estados pero son
    jurisdicciones bajo soberanía estadounidense
63. #   es decir, son ciudadanos estadounidenses, reciben servicios federales como
    salud pública y vacunas y participan en programas como VAERS

64. #Revisamos de nuevo los valores únicos de las columnas categóricas
65. categoricas = [
66. "SEX", "DIED", "L_THREAT", "ER_VISIT", "HOSPITAL",
67. "DISABLE", "RECOVD", "BIRTH_DEFECT", "STATE"
68. ]

69. for col in categoricas:
70. print(f"\n🔍 Valores únicos en la columna '{col}':")
71. print(df2_VAERS_2021_Data[col].unique())

72. #Eliminamos las filas con NaN en la columna STATE, ya que no tienen sentido
73. df2_VAERS_2021_Data = df2_VAERS_2021_Data.dropna(subset=["STATE"])

74. df2_VAERS_2021_Data.shape # Pasamos de (768386, 17) a (651494, 17)

75. df2_VAERS_2021_Data.isnull().mean().sort_values(ascending=False)

76. #Veo que hay muchos valores con NaN, pero tiene su explicación, ya que estas
    columnas significan cosas como:
77. #   Defectos en nacimiento, fecha de muerte, muerte, amenaza vital, etc. Y si
    hay un NaN implica que no se ha reportado, no que no se ha producido
78. #   Es decir, que podemos sustituir algunas de estas columnas los NaN por 'No'

79. #Vamos a hacer esto con las columnas que claramente si es NaN se puede poner que
    No:
80. columnas_binarias = ["DIED", "L_THREAT", "HOSPITAL", "DISABLE", "BIRTH_DEFECT",
    "ER_VISIT"]

81. for col in columnas_binarias:
82. df2_VAERS_2021_Data.loc[:, col] = df2_VAERS_2021_Data[col].fillna("N")

83. #Perfecto, ahora mucho mejor. Datedied y Hospdays no puedo ponerles que 'No'
    porque son fechas o días.
84. df2_VAERS_2021_Data.isnull().mean().sort_values(ascending=False)

85. #Ahora vamos a revisar cuales son los rangos de fechas que se manejan en el
    dataset para ver inconsistencias

```

```

86. #Columnas a analizar:
87. date_cols = ["RECVDATE", "VAX_DATE", "ONSET_DATE", "DATEDIED"]

88. # Mostrar rango de fechas por cada columna
89. for col in date_cols:
90.     print(f"📅 {col} → Desde: {df2_VAERS_2021_Data[col].min()} hasta:
      {df2_VAERS_2021_Data[col].max()}")

91. #Vemos que hay columnas con fechas sin sentido, seguramente de fallos a la hora
    de pasar datos y tal. Vamos a poner fecha min y max

92. # Pongo fechas un poco por encima y por abajo por si acaso, aunque la vacunación
    empezó oficialmente en EEUU el 14 de diciembre con Pfizer. Fecha máxima 2022
    porque el dataset es de 2021
93. fecha_min = pd.to_datetime("2020-12-01")
94. fecha_max = pd.to_datetime("2022-01-15")
95. # Vamos a filtrar las fechas para quedarnos solo con las que están dentro de
    este rango
96. df2_VAERS_2021_Data = df2_VAERS_2021_Data[
97.     (df2_VAERS_2021_Data["VAX_DATE"].isna() | ((df2_VAERS_2021_Data["VAX_DATE"] >=
        fecha_min) & (df2_VAERS_2021_Data["VAX_DATE"] <= fecha_max))) &
98.     (df2_VAERS_2021_Data["ONSET_DATE"].isna() | ((df2_VAERS_2021_Data["ONSET_DATE"]
        >= fecha_min) & (df2_VAERS_2021_Data["ONSET_DATE"] <= fecha_max))) &
99.     (df2_VAERS_2021_Data["DATEDIED"].isna() | ((df2_VAERS_2021_Data["DATEDIED"] >=
        fecha_min) & (df2_VAERS_2021_Data["DATEDIED"] <= fecha_max)))
100. ]

101. # Volvemos a revisar las columnas
102. date_cols = ["RECVDATE", "VAX_DATE", "ONSET_DATE", "DATEDIED"]

103. # Mostrar rango de fechas por cada columna
104. for col in date_cols:
105.     print(f"📅 {col} → Desde: {df2_VAERS_2021_Data[col].min()} hasta:
      {df2_VAERS_2021_Data[col].max()}")

106. df2_VAERS_2021_Data.shape # Pasamos de (651494, 17) a (646175, 17), eliminando
    datos de fechas incorrectos.

107. df2_VAERS_2021_Data.describe() #He pensado que quizás habría valores numéricos
    mal
108. # Todo está bien salvo HOSPDAYS, nadie se pega 99999 días en el hospital
109. # Voy a comprobar quien ha estado más de un año en un hospital, las eliminaré y
    listo

110. df2_VAERS_2021_Data.shape

111. # Listo, elimino los que tienen mas de un año porque es o muy raro o falso. Hay
    que poner los NaN también ya que significan que no han estado en el hospital
112. df3_VAERS_2021_Data = df2_VAERS_2021_Data[

```

```

113. (df2_VAERS_2021_Data["HOSPDAYS"] <= 365) |
      (df2_VAERS_2021_Data["HOSPDAYS"].isna())
114. ]

115. df3_VAERS_2021_Data.shape

116. df3_VAERS_2021_Data.describe()

117. # Con esto ya hemos acabado el analisis exploratorio y podemos empezar en Power
      BI con este dataset. Repasemos lo hecho o no hecho aquí:
118. # ✅ 2.1) Limpiar los NaN
119. # ✅ 2.2) Formato datetime
120. # ⌚ 2.3) Modificar decimales --> En PBI
121. # ✅ 2.4) Revisar duplicados
122. # ✅ 2.5) Revisar columnas categóricas
123. # ✅ 2.6) Outliers (IQR, boxplots) --> Outliers y medidas extrañas comprobados
      con describe()
124. # ❌ 2.7) Distribución de Datos (histogramas, frecuencias) --> Idem, en PBI
125. # ✅ 2.8) Filtrar fechas previas a la pandemia
126. # ❌ 2.9) Correlaciones --> PBI
127. # ✅ 2.10) Eliminar columnas innecesarias
128. # ✅ 2.11) Revisar valores de desviacion como minimos y maximos de columnas

129. # Con esto ya hemos acabado el analisis exploratorio y podemos empezar en Power
      BI con este dataset.
130. df3_VAERS_2021_Data.to_csv("df3VAERS2021Data_limpio_para_powerbi.csv",
      index=False)

```

4

Cleaning of VAERS Database *symptoms.csv*

```

1. # Author: Miguel Atarés
2. # coding: utf-8
3. # Análisis de la Base de Datos: SYMPTOMS2021

4. df_SYMPTOMS_2021_Data = pd.read_csv(r'C:\Users\migue\Desktop\COVID-19
      Project\Databases\2021VAERSData\2021VAERSSYMPTOMS.csv', encoding = "ISO-8859-1",
      delimiter=',')

5. #Primero eliminamos las columnas:
6. # Lista de columnas que NO se van a conservar
7. columnas_a_eliminar = [
8.     "SYMPTOMVERSION1",
9.     "SYMPTOMVERSION2",
10.    "SYMPTOMVERSION3",
11.    "SYMPTOMVERSION4",
12.    "SYMPTOMVERSION5",
13.    "ORDER"

```

```

14. ]

15. # Eliminar esas columnas del DataFrame
16. df2_SYMPTOMS_2021_Data = df_SYMPTOMS_2021_Data.drop(columns=columnas_a_eliminar)

17. df2_SYMPTOMS_2021_Data.shape

18. df2_SYMPTOMS_2021_Data.head()

19. #Claro, no hay más variables numéricas realmente
20. df2_SYMPTOMS_2021_Data.describe()

21. #Tipo de datos está ok
22. df2_SYMPTOMS_2021_Data.info()

23. #limpiamos duplicados
24. df2_SYMPTOMS_2021_Data.drop_duplicates(inplace=True)

25. df2_SYMPTOMS_2021_Data.shape #Pasamos de (1030218, 6) a (1028726, 6) por lo que
    había algún duplicado

26. #Siempre y cuando esté el ID y el SYMPTOM1 todo está ok.
27. df2_SYMPTOMS_2021_Data.isnull().mean().sort_values(ascending=False)

28. # Con esto ya hemos acabado el analisis exploratorio y podemos empezar en Power
    BI con este dataset. Repasemos lo hecho o no hecho aquí:
29. # ☒ 2.1) Limpiar los NaN --> Los que hay es normal que los haya
30. # ☒ 2.2) Formato datetime --> No hay
31. # ☒ 2.3) Modificar decimales --> No hay
32. # ☒ 2.4) Revisar duplicados
33. # ☒ 2.5) Revisar columnas categóricas --> No hay
34. # ☒ 2.6) Outliers (IQR, boxplots) --> No hay
35. # ☒ 2.7) Distribución de Datos (histogramas, frecuencias) --> PBI
36. # ☒ 2.8) Filtrar fechas previas a la pandemia --> No hay
37. # ☒ 2.9) Correlaciones --> PBI
38. # ☒ 2.10) Eliminar columnas innecesarias
39. # ☒ 2.11) Revisar valores de desviacion como minimos y maximos de columnas -->
    No hay

40. # Con esto ya hemos acabado el analisis exploratorio y podemos empezar en Power
    BI con este dataset.
41. df2_SYMPTOMS_2021_Data.to_csv("df2VAERS2021Symptoms_limpio_para_powerbi.csv",
    index=False)

```

Cleaning of VAERS Database VAX.csv

```

1. # Author: Miguel Atarés
2. # coding: utf-8
3. # Análisis de la Base de Datos: SYMPTOMS2021

4. df_SYMPTOMS_2021_Data = pd.read_csv(r'C:\Users\miguel\Desktop\COVID-19
Project\Databases\2021VAERSData\2021VAERSSYMPTOMS.csv', encoding = "ISO-8859-1",
delimiter=',')

5. #Primero eliminamos las columnas:
6. # Lista de columnas que NO se van a conservar
7. columnas_a_eliminar = [
8. "SYMPTOMVERSION1",
9. "SYMPTOMVERSION2",
10. "SYMPTOMVERSION3",
11. "SYMPTOMVERSION4",
12. "SYMPTOMVERSION5",
13. "ORDER"
14. ]

15. # Eliminar esas columnas del DataFrame
16. df2_SYMPTOMS_2021_Data = df_SYMPTOMS_2021_Data.drop(columns=columnas_a_eliminar)

17. df2_SYMPTOMS_2021_Data.shape

18. df2_SYMPTOMS_2021_Data.head()

19. #Claro, no hay más variables numéricas realmente
20. df2_SYMPTOMS_2021_Data.describe()

21. #Tipo de datos está ok
22. df2_SYMPTOMS_2021_Data.info()

23. #limpiamos duplicados
24. df2_SYMPTOMS_2021_Data.drop_duplicates(inplace=True)

25. df2_SYMPTOMS_2021_Data.shape #Pasamos de (1030218, 6) a (1028726, 6) por lo que
había algún duplicado

26. #Siempre y cuando esté el ID y el SYMPTOM1 todo está ok.
27. df2_SYMPTOMS_2021_Data.isnull().mean().sort_values(ascending=False)

28. # Con esto ya hemos acabado el analisis exploratorio y podemos empezar en Power
BI con este dataset. Repasemos lo hecho o no hecho aquí:
29. # ✅ 2.1) Limpiar los NaN --> Los que hay es normal que los haya
30. # ✅ 2.2) Formato datetime --> No hay

```

```

31. # ☒ 2.3) Modificar decimales --> No hay
32. # ☒ 2.4) Revisar duplicados
33. # ☒ 2.5) Revisar columnas categóricas --> No hay
34. # ☒ 2.6) Outliers (IQR, boxplots) --> No hay
35. # ☒ 2.7) Distribución de Datos (histogramas, frecuencias) --> PBI
36. # ☒ 2.8) Filtrar fechas previas a la pandemia --> No hay
37. # ☒ 2.9) Correlaciones --> PBI
38. # ☒ 2.10) Eliminar columnas innecesarias
39. # ☒ 2.11) Revisar valores de desviación como mínimos y máximos de columnas -->
    No hay

40. # Con esto ya hemos acabado el análisis exploratorio y podemos empezar en Power
    BI con este dataset.
41. df2_SYMPTOMS_2021_Data.to_csv("df2VAERS2021Symptoms_limpio_para_powerbi.csv",
    index=False)

42. #   Análisis de la Base de Datos: VAX2021
43. # </h1>

44. df_VAX_2021_Data = pd.read_csv(r'C:\Users\migue\Desktop\COVID-19
    Project\Databases\2021VAERSData\2021VAERSVAX.csv', encoding = "ISO-8859-1",
    delimiter=',')

45. # Eliminamos columnas que no nos interesan
46. columnas_a_eliminar_VAX = [
47.     "VAX_LOT",
48.     "VAX_ROUTE",
49.     "VAX_SITE",
50.     "VAX_NAME",
51.     "ORDER"
52. ]

53. # Eliminar columnas no necesarias
54. df2_VAX_2021_Data = df_VAX_2021_Data.drop(columns=columnas_a_eliminar_VAX)

55. df2_VAX_2021_Data.head()

56. #Quizás en PowerBI las interprete como dosis de refuerzo a partir de la tercera.
57. df2_VAX_2021_Data['VAX_DOSE_SERIES'].unique()

58. # Solo necesito las de COVID
59. df2_VAX_2021_Data['VAX_TYPE'].unique()

60. # Todo lo que contenga la palabra "COVID19" se elimina de la columna VAX_TYPE
61. df3_VAX_2021_Data =
    df2_VAX_2021_Data[df2_VAX_2021_Data["VAX_TYPE"].str.contains("COVID19",
    case=False, na=False)]

```

```

62. #Perfecto, ahora tenemos en nuestro database solo estos valores, pero aun así
    ahora voy a eliminar la columna porque es redundante
63. df3_VAX_2021_Data['VAX_TYPE'].unique()

64. #Elimino columna VAX_TYPE
65. df4_VAX_2021_Data = df3_VAX_2021_Data.drop(columns=["VAX_TYPE"])

66. df4_VAX_2021_Data.shape

67. #Eliminamos duplicados
68. df4_VAX_2021_Data.drop_duplicates(inplace=True)

69. df4_VAX_2021_Data.shape #Pasamos de (757765, 3) a (742210, 3)

70. #Miro tipo de datos. Sinceramente creo que las dosis deberían ser un entero,
    igual elimino los UNK.
71. df4_VAX_2021_Data.info()

72. #La verdad que hay pocos NaN, solo un 0,3%.
73. df4_VAX_2021_Data.isnull().mean().sort_values(ascending=False)

74. #Poca categorización aquí de tipos de vacunas, pero bueno, mantenemos.
75. df4_VAX_2021_Data['VAX_MANU'].unique()

76. # Vamos a ver el % de 'UNK' en "VAX_DOSE_SERIES" a ver si podemos eliminarlo.
77. total = len(df4_VAX_2021_Data)

78. # Filas con 'UNK'
79. unk_count = (df4_VAX_2021_Data["VAX_DOSE_SERIES"] == "UNK").sum()
80. unk_pct = round((unk_count / total) * 100, 2)

81. print(f"'UNK': {unk_count} filas ({unk_pct}%)")

82. # Tiene un 15,36% que es mucho, por lo que creo que lo voy a dejar en el
    análisis.
83. # Si veo que me molesta mucho en power BI a la hora de hacer cuentas, quizás
    copio y pego la columna y la quito y listo.

84. # Con esto ya hemos acabado el analisis exploratorio y podemos empezar en Power
    BI con este dataset. Repasemos lo hecho o no hecho aquí:
85. # ☒ 2.1) Limpiar los NaN --> Los asumimos
86. # ☒ 2.2) Formato datetime --> No hay
87. # ☒ 2.3) Modificar decimales --> No hay
88. # ☒ 2.4) Revisar duplicados
89. # ☒ 2.5) Revisar columnas categóricas
90. # ☒ 2.6) Outliers (IQR, boxplots) --> No hay
91. # ☒ 2.7) Distribución de Datos (histogramas, frecuencias) --> PBI
92. # ☒ 2.8) Filtrar fechas previas a la pandemia --> No hay
93. # ☒ 2.9) Correlaciones --> PBI

```

