

6 Annexes.

6.1 Sequence realign.

In this annex the codes used to realign the sequences taken from [1]. These sequences are aligned following the KABAT scheme and have to be realigned to follow the AHO one.

```
1 from Bio.Seq import Seq
2 from Bio import SeqIO
3 from Bio.SeqRecord import SeqRecord
4 import subprocess as sub
5
6 seqs_list = list(SeqIO.parse(open('/home/david/Desktop/TFM/DatosAstati/ModAsti/
    allYES_uniqueAA_c90-140_IGHV1-2_IGHJ2_SEQMA45_3SE9GG_reH2_allSEQMA0.afasta'),
    fasta'))
7 # seqs_list = list(SeqIO.parse(open('/home/david/Desktop/TFM/DatosAstati/ModAsti/
    S2_cluVJ_reH2_allSEQMA0.afasta'), 'fasta'))
8
9 def quita_gaps(seq):
10     cand = ''
11     for amin in seq:
12         if amin != '-':
13             cand += amin
14     return cand
15
16 def Ident(seqs_list):
17     ident = []
18     for i in range(len(seqs_list)):
19         ident.append(seqs_list[i].id)
20     return ident
21
22 def escribe_fasta(Chains, Ident, f_out):
23     for N_C in range(len(Chains)):
24         record = SeqRecord(Seq(str(Chains[N_C])), id=Ident[N_C], description='Asti seq
            no align')
25         SeqIO.write(record, f_out, 'fasta')
26
27 def desalinea(seqs_list):
28     d_out = 'No_align.fasta'
29     f_out = open(d_out, 'w')
30     seqs = []
31     ids = Ident(seqs_list)
32     for j in range(len(seqs_list)):
33         seqs.append(quita_gaps(seqs_list[j].seq))
34     escribe_fasta(seqs, ids, f_out)
35     return d_out
36
37 def alinea(f_in):
38     output = 'Alineados.txt'
39     sub.run(['ANARCI', '-i', f_in, '-o', output, '-r', 'heavy', '-s', 'a', '--
        use_species', 'human',
40         '--ncpu', '12'], capture_output = True)
41     sub.run(['rm', f_in])
42     return(output)
43
44 def ANARCI_fasta(align_in):
45     with open(align_in, 'r') as file:
46         Seqs = []
```

```

47     Iden = []
48     imp_iden = []
49     seq = ''
50     for line in file:
51         splitted_line = line.split()
52         if len(splitted_line) == 6:
53             prov_iden = splitted_line[1]
54             if splitted_line[0] == 'H':
55                 seq += splitted_line[2]
56             if splitted_line[0] == '//':
57                 if len(seq) == 0 : imp_iden.append(prov_iden)
58                 if len(seq) <= 149 and len(seq) > 0:
59                     seq = seq + (149 - len(seq))*'-'
60                     Seqs.append(seq)
61                     Iden.append(prov_iden)
62             seq = ''
63     sub.run(['rm', align_in])
64     f_out = open('Asti_Aligned.fasta','w')
65     print(len(Iden))
66     print(len(imp_iden))
67     escribe_fasta(Seqs, Iden, f_out)
68     return imp_iden
69
70 def extrae_data(ident_imp, ref_data):
71     ident_tot = []
72     final_items = []
73
74     for item in ref_data:
75         ident_tot.append(item.id)
76
77     for id_c in ident_imp:
78         ind = ident_tot.index(id_c)
79         final_items.append(ref_data[ind])
80     return final_items
81
82 # seqs_list = list(SeqIO.parse(open('./germ_clust.fasta'),'fasta'))
83 # seqs_list = list(SeqIO.parse(open('./hiper_mut_clust.fasta'),'fasta'))
84 seqs_list = list(SeqIO.parse(open('../Extrae_cluster_raro/filtradas_all_yes.fasta'),
85                               , 'fasta'))
86
87 print('Seqs leidas:', len(seqs_list))
88
89 f_in = desalineada(seqs_list)
90 print('Seqs desalineadas')
91
92 txt_align = alineada(f_in)
93 print('Seqs realineadas')
94
95 imp_ident = ANARCI_fasta(txt_align)
96
97 imp_final = extrae_data(imp_ident, seqs_list)
98 print('Impossible to align: ', len(imp_final))
99
100 # file_imp = 'imp_to_align.fasta'
101 # with open('hiper_mut_clust.fasta', 'w') as file:
102 #     SeqIO.write(imp_final, file, 'fasta')

```

6.2 Sequence generation.

In this annex the codes used to generate the different types of data bases will be presented.

6.2.1 Gaussian data bases.

To generate the data sets of Gaussian sequences there are two possible options depending on how we want to generate them. If you want to generate a large number of sequences at once or if you want to generate small groups of sequences. If you are in the first case, you should use the python code and if you are in the second case, the julia code. This is because the time it takes Julia to generate sequences grows linearly with the number of sequences, but in the case of python it is not like that, it takes much more time than Julia to generate a small number of sequences, but as the time it takes does not scale linearly with the number of sequences, it ends up being much faster generating large blocks of sequences.

The Python script:

```
1 import numpy as np
2 import sys
3 from scipy.stats import multivariate_normal
4 from Bio.Seq import Seq
5 from Bio import SeqIO
6 from Bio.SeqRecord import SeqRecord
7
8 sys.path.append("/home/david/Desktop/TFM/seqsTFM-base/seqsTFM/")
9 from transform import *
10
11 arg = len(sys.argv)
12 if len(sys.argv) == 2:
13     for i in range(len(sys.argv)):
14         Ncad = int(sys.argv[1])
15 else:
16     print('Num de args incorrecto.')
17     print('Usage: ')
18     print('GenSem.py Ncad')
19     exit()
20
21 print('Secuencias a generar: ', Ncad)
22
23
24 ##### Funciones #####
25
26 def gen_semilla(cov, mean, Nsem, Limite, Lcad):
27     '''
28     A partir de la matriz inversa de la covarianza y de la media, genera una
29     una distribucion gaussiana multivariada con la cual genera las semillas.
30     '''
31     cadbin = np.random.multivariate_normal(mean, cov, Nsem)
32     cadascii = []
33     ident = []
34     auxascii = np.zeros((Nsem, Lcad), dtype = np.uint8)
35     for j in range(Nsem):
36         auxascii[j] = binary_to_uint8(floattoint(cadbin[j], Limite))
37         cadascii.append("".join([(chr(a)) for a in auxascii[j]]))
```

```

38     ident.append(str(j))
39     return cadascii, ident
40
41 ##### Parametros #####
42
43 Limite = 0.367
44 Lcad = 298
45
46 ##### Main #####
47
48 ref = np.loadtxt('./Modelo/MedBin.txt')
49 ref = np.loadtxt('./Modelo/CovBin.txt')
50 print('Modelo cargado')
51
52 cadascii, ident = gen_semilla(cov = ref, mean = ref, Nsem = Ncad, Limite =
    Limite, Lcad = Lcad)
53 print('Semilla generada')
54
55 with open('Gauss_seqs_' + str(Ncad) + '.fasta', 'w') as file:
56     escribe_fasta(cadascii, ident, file)

```

The Julia script:

```

1 using LinearAlgebra
2 using Random
3 using Distributions
4
5 function gen_obj_gauss( ref, ref)
6
7     cov = ( ref + ref')/2
8     Dist = MvNormal( ref, cov)
9     return (Dist)
10 end
11
12 function gen_delt_sem(delt_N, LcadB, Dist)
13
14     CadenaP = Array{Float64}(undef, LcadB)
15     Cadena = Matrix{Float64}(undef, delt_N, LcadB)
16     for i in 1:delt_N
17         rand!(Dist, CadenaP)
18         Cadena[i, :] = CadenaP
19     end
20     return Cadena
21 end

```

6.2.2 IGoR data bases.

To generate the IGoR databases, as already mentioned in the Section 2.4 section, there are several steps and therefore several programs are needed. The first of them is the main script which is supported by the two scripts that will be presented later. In this main script are the necessary functions to not only generate the nucleotide sequences with IGoR, but also align them with ANARCI.

```

1 import subprocess as sub
2 import numpy as np
3 import csv

```

```

4 from Bio import SeqIO
5
6 def genera_seqs():
7     sub.run(['igor', '-species', 'human', '-chain', 'heavy_naive'], capture_output = True
8     )
9     sub.run(['igor', '-generate', str(len_bloque), '-set_wd', './'], capture_output =
10     True )
11
12 def comprueba():
13     with open('./generated/generated_seqs_werr.csv', 'r') as file:
14         reader = csv.reader(file)
15         a = int(len(list(reader)))
16         if a > len_bloque:
17             return (True)
18         else:
19             return (False)
20
21 def csv_to_fasta(i):
22     output = 'Igor_S' + str(i) + '.fasta'
23     sub.run(['python3', 'Igor_Trl.py', './generated/generated_seqs_werr.csv', output],
24     capture_output = True)
25     return(output)
26
27 def alineaa(f_in, i):
28     output = 'Igor_Aling' + str(i) + '.txt'
29     sub.run(['ANARCI', '-i', f_in, '-o', output, '-r', 'heavy', '-s', 'a',
30     '--ncpu', '12', '--use_species', 'human'], capture_output = True)
31     sub.run(['rm', f_in])
32     return(output)
33
34 def ANARCI_fasta(aligned_in, i):
35     output = 'Igor_Aling' + str(i) + '.fasta'
36     sub.run(['python3', 'ANARCI_txt.py',
37     aligned_in, output], capture_output = True)
38     sub.run(['rm', aligned_in])
39
40 long_base = 2000000
41 len_bloque = 10000
42 n_files = range(0, int(long_base/len_bloque) + 1)
43 print('Archivos a generar: ', int(long_base/len_bloque) + 1)
44
45 for i in (n_files):
46     while True:
47         genera_seqs()
48         if comprueba() == True:
49             break
50     f_out = csv_to_fasta(i)
51     aligned_out = alineaa(f_out, i)
52     ANARCI_fasta(aligned_out, i)
53     print('Avance: ', np.round((i/len(n_files))*100), '%', end = '\r')
54
55 with open("/home/david/Desktop/TFM/Igor_Gen_Data/Igor_50k.fasta", "a") as final:
56     for j in n_files:
57         try:
58             file_in = 'Igor_Aling' + str(j) + '.fasta'
59             record = SeqIO.parse(file_in, 'fasta')
60             SeqIO.write(record, final, 'fasta')

```

```

60     sub.run(['rm', file_in])
61 except:
62     print('File ', j, 'not found.')

```

In the other two scripts, the sequences are translated and changes are made to their format so one program can correctly receive the output of another. Specifically, the first program receives as input the csv that has IGoR as output, translates the sequences and converts them into fasta format, which is the output format. The second program only makes a format change: from the txt type that has ANARCI as output to fasta format to store the sequences correctly.

```

1  #A partir de un archivo CSV salido de Igor construye un archivo fasta con las
   secuencias ya traducidas.
2
3  from Bio.Seq import Seq
4  from Bio import SeqIO
5  from Bio.SeqRecord import SeqRecord
6  import pandas as pd
7  import numpy as np
8  import csv
9  import sys
10
11 N_Arg = len(sys.argv)
12 if N_Arg == 3:
13     File_Ini = sys.argv[1]
14     File_Output = sys.argv[2]
15 else:
16     print('Igor_Trl.py Seq_File Input Output')
17     print('Seq_File: file which contains the sequences.')
18     print('Output File.')
19     exit()
20
21 def Escribe_Fasta(Chains,Ident,File):
22     ''' A partir de dos listas, una con las cadenas y otra con sus identificadores,
23     crea un archivo fasta'''
24     for N_C in range(len(Chains)):
25         record = SeqRecord(Seq(str(Chains[N_C])),id=Ident[N_C],
26                             description="Generated by Igor")
27         #Une las listas de cadenas e identificadores para pasarlas al formato de dato
28         #que usa biopython y poder usar con ellas sus funciones.
29         SeqIO.write(record,File,'fasta')
30
31
32 def CSV_to_Fasta(File_Ini):
33     '''A partir de un archivo CSV con las secuencias crea uno fasta'''
34     Chains_A = []
35     Ident_A = []
36     #Listas donde almacenamos las secuencias de ARN y sus identificadores.
37     with open(File_Ini, 'r') as file:
38         #F.csv es un archivo CSV con las cadenas de Igor.
39         reader = csv.reader(file)
40         next(reader, None)
41         #Saltamos la cabecera que le pone Igor a los archivos de secuencias.
42         for row in reader:
43             row = row[0].split(";")
44             Ident_A.append(Genera_Iden_ARN(row[0],row[1]))
45             #La cabecera ser un contador y la longitud de la secuencia.
46             Chains_A.append(row[1])

```

```

47     return(Chains_A, Ident_A)
48
49 def Genera_Iden_ARN(Cont, Chain):
50     '''
51     Will create every chain identification.
52     '''
53     A_Len = len(Chain)
54     #Longitud de la cadena.
55     if A_Len % 3 == 0:
56         Prod = 'P'
57     else:
58         Prod = 'NP'
59     Iden = Cont+'|'+str(A_Len)+'|'+str(Prod)
60     return(Iden)
61     #La cabecera ser un contador y la longitud de la secuencia.
62
63 def Translate_ARN(Chains_A, Ident_A):
64     '''A partir de las cadenas ARN produce un archivo fasta con las cadenas
65     de aminoacidos'''
66     Chains_P = []
67     Ident_P = []
68     #Listas donde almacenamos las secuencias de proteínas y sus identificadores.
69     Cont = 0
70     #Util para las cabeceras de las secuencias.
71     for Chain in Chains_A:
72         P_len = len(Chain)
73         #Longitud de la cadena.
74         Chains_P.append(Seq(Chain).translate(to_stop=True))
75         #La cabecera ser un contador y la longitud de la secuencia.
76         Ident_P.append(str(Cont)+'|'+str(P_len))
77         Cont += 1
78     Prot_F = open(File_Output, 'w')
79     Escribe_Fasta(Chains_P, Ident_P, Prot_F)
80
81 Chains, Ident = CSV_to_Fasta(File_Ini)
82 Translate_ARN(Chains, Ident)

```



```

1 from Bio.Seq import Seq
2 from Bio import SeqIO
3 from Bio.SeqRecord import SeqRecord
4 import pandas as pd
5 import numpy as np
6 import csv
7 import sys
8
9
10 N_Arg = len(sys.argv)
11 if N_Arg == 3:
12     File_Ini = sys.argv[1]
13     Output = sys.argv[2]
14 else:
15     print('Trans_Def.py Seq_File \t Chain_File_Type \t Output_File')
16     print('Seq_File: file which contains the sequences.')
17     exit()
18
19 def Escribe_Fasta(Chains, Ident, Output):
20     File_0 = open(Output, 'w')
21     for N_C in range(len(Chains)):
22         record = SeqRecord(Seq(str(Chains[N_C])), id=Ident[N_C],

```

```

23     description="Generated by Igor")
24     #Do not change the description if then Igor_Trl.py is going to be launched.
25     SeqIO.write(record,File_0,'fasta')
26
27 with open(File_Ini, 'r') as file:
28     Seqs = []
29     Iden = []
30     Cont = 0
31     Esc = 0
32     P = 0
33     for line in file:
34         split_line = line.split()
35         if len(split_line)>2 and split_line[2] == 'Generated':
36             Marca = Cont + 5
37             P_Iden = split_line[0] + split_line[1]
38             P_Chain = ''
39             if Cont == Marca and line.split('|')[1] == 'human':
40                 Esc = 1
41                 P += 1
42             if split_line[0] == 'H':
43                 P_Chain += split_line[2]
44             if split_line[0] == '//':
45                 if Esc == 1 and len(P_Chain) <= 149:
46                     Seqs.append(P_Chain + (149 - len(P_Chain))*'-')
47                     Iden.append(P_Iden)
48                     Esc = 0
49             Cont += 1
50 print(len(Seqs),'\t', P)
51 Escribe_Fasta(Seqs,Iden,Output)

```

6.3 Null columns calculation.

In this annex is presented a Python script with the following functionalities:

- Columns which contains all the possible amino acids calculation.
- Null columns calculation.
- Comparative between the null columns of a reference data set and another data set.
- Variability column calculation.

```

1 import numpy as np
2 import numba as nb
3 import sys
4 from Bio import SeqIO
5 import matplotlib.pyplot as plt
6
7 sys.path.append("../seqsTFM-base/seqsTFM/")
8 from transform import ascii_to_binary_vect, ref_seqs_mat
9
10 amin_dict_8 = np.asarray([65, 67, 68, 69, 70, 71, 72, 73, 75, 76, 77, 78, 80, 81,
11                          82, 83, 84, 86, 87, 89])
12 ##### Functions #####
13

```

```

14 def all_amin_ascii(asti_mat):
15     '''
16     Returns the number of columns containing all amino acids.
17     '''
18     ncad, lcad = asti_mat
19     asti_mat_T = asti_mat.T
20     recuento = np.zeros(lcad)
21     for pos in range(lcad):
22         if check_all_amin(asti_mat_T[pos]) == len(amin_dict_8):
23             recuento[pos] = 1
24     return recuento
25
26 def check_all_amin(asti_mat_T_i):
27     '''
28     Check if a column contains all amino acids.
29     '''
30     count = 0
31     for amin in amin_dict_8:
32         if amin not in asti_mat_T_i:
33             break
34         count += 1
35     return count
36
37 def all_col_0(seqs_mat_B):
38     '''
39     Returns an array containig the null columns.
40     '''
41     nseq, lcad_B = seqs_mat_B.shape
42     seqs_mat_B = seqs_mat_B.T
43     filas_0 = []
44     for fila in range(lcad_B):
45         if check_col(seqs_mat_B[fila], nseq) == nseq:
46             filas_0.append(fila)
47     return np.asarray(filas_0)
48
49 @nb.njit(parallel = False, fastmath = True)
50 def check_col(col, nseq):
51     '''
52     Counts the number of 0's in a column.
53     '''
54     count = 0
55     for i in range(nseq):
56         if col[i] == 0:
57             count += 1
58     return count
59
60 def no_cont(filas_ref, filas_p):
61     '''
62     Returns a vector with the rows that are null in rows_p, but not null in rows_ref.
63     '''
64     no_cont = []
65     for i in range(len(filas_p)):
66         cont = False
67         for j in range(len(filas_ref)):
68             if int(filas_p[i]) == int(filas_ref[j]):
69                 cont = True
70                 break
71         if cont == False:
72             no_cont.append(filas_p[i])

```

```

73     return np.asarray(no_cont)
74
75 def comp_varia(ref_mat_B, col):
76     '''
77     Returns the variability of the columns contained in a vector.
78     '''
79     nseq, lcad_B = ref_mat_B.shape
80     ref_mat_B = ref_mat_B.T
81     cons = np.zeros(len(col))
82     for i in range(len(col)):
83         cons[i] = check_col(ref_mat_B[col[i]], nseq)/nseq
84     return cons
85
86 def plot_var(no_comun, cons, nseqs_ref):
87     plt.scatter(no_comun, 1 - cons)
88     plt.axhline(y = 1/nseqs_ref, color='r', linestyle=':')
89     plt.xlabel('No coincident columns')
90     plt.ylabel('Variability in ref. seqs.')
91     plt.show()
92
93 def save_data(no_comun, cons, data_file):
94     data = np.zeros((2, len(no_comun)))
95     data[0] = no_comun
96     data[1] = cons
97     data_out = data_file.split('.')[0] + '.txt'
98     with open('./Data/Var/Var_' + data_out, 'w') as out:
99         np.savetxt(out, data)
100
101 ##### Main code #####
102
103 #### Data uploading ####
104 data_file = 'Igor_1M_798k.fasta'
105 seqs = list(SeqIO.parse(open('../DataBases/Igor/' + data_file), 'fasta'))
106 seqs_mat_B = ascii_to_binary_vect(ref_seqs_mat(seqs))
107 print('Seq leidas:', len(seqs))
108
109 ref_seqs = list(SeqIO.parse(open('../DataBases/Bases_Asti/S1yS2_realign/
    Asti_Aligned_S1.fasta'), 'fasta'))
110 ref_mat_B = ascii_to_binary_vect(ref_seqs_mat(ref_seqs))
111 print('Seq ref:', len(ref_seqs))
112
113 #### Calculation of null rows ####
114 filas_0_seqs = all_col_0(seqs_mat_B)
115 print('Filas nulas en seqs:', len(filas_0_seqs))
116
117 filas_0_ref = all_col_0(ref_mat_B)
118 print('Filas nulas en ref:', len(filas_0_ref))
119
120 #### Common columns ####
121 no_comun = no_cont(filas_ref = filas_0_ref, filas_p = filas_0_seqs)
122 print('Filas nulas no comunes:', len(no_comun))
123
124 cons = comp_varia(ref_mat_B, no_comun)
125 # plot_var(no_comun, cons, len(ref_seqs))
126 save_data(no_comun, cons, data_file)

```

6.4 Large data bases covariance and mean calculation.

In this annex is presented a Python script that contains a class object that is able to calculate the mean and the covariance of a large data set. It can be done also considering a prior distribution for the sequences.

```
1 import numpy as np
2 import numba as nb
3 import sys
4
5 sys.path.append("./")
6 from stats import posterior
7
8 ##### Classes #####
9 class System:
10     '''
11     In this class we will store both the covariance and the mean of a data set.
12     of a data set. The key is that it will allow us to update both as we add
13     sequences.
14     This way the size of the data set is not limited by RAM memory.
15     '''
16     def __init__(self, Lcad_B):
17         '''
18         Initial parameters: mean and cov, both initialized to 0 and the binary length
19         of the sequences we are going to use.
20         '''
21         self.cov = np.zeros((Lcad_B, Lcad_B))
22         self.mean_t = np.zeros(Lcad_B)
23         self.Lcad_B = Lcad_B
24
25     def update_no_prior(self, seqs_N, n_seq_i, N):
26         '''
27         seqs_N: sequences to be added to the data set.
28         n_seq_i: number of sequences considered in the data set so far.
29         N: number of sequences to add to the data set in this step.
30         From these 3 things it updates the value of the mean and cov assuming
31         a value of = 0.
32         '''
33         old_weight, N_weight = calc_weight(n_seq_i, N)
34         mean_N, cov_intra_N = self.cov_and_mean_N(seqs_N)
35         self.update_cov_and_mean(old_weight, N_weight, cov_intra_N, mean_N)
36
37     def update(self, seqs_N, n_seq_i, N, old, oldU, ):
38         '''
39         seqs_N: sequences to be added to the data set.
40         n_seq_i: number of sequences considered in the data set so far.
41         N: number of sequences to add to the data set in this step.
42         old, oldU: a priori distributions for the mean and cov, respectively.
43         : fraction of prior to include.
44         From these 5 things it updates the value of the mean and cov.
45         '''
46         old_weight, N_weight = calc_weight(n_seq_i, N)
47         mean_N, cov_intra_N = self.cov_and_mean_N(seqs_N)
48         Post_mean_N, Post_cov_intra_N = posterior(mean_N, cov_intra_N, old, oldU, self.Lcad_B)
49         self.update_cov_and_mean(old_weight, N_weight, Post_cov_intra_N, Post_mean_N)
```

```

50 def update_cov_and_mean(self, weight_old, weight_N, cov_intra_N, mean_N):
51     '''
52     Update the mean and covariance from the weights and the intra-group mean and
53     covariance of the sequences to be added.
54     '''
55     self.mean_t = self.mean_t * weight_old + mean_N * weight_N
56     dif = (mean_N - self.mean_t).reshape(self.Lcad_B, 1)
57     cov_inter_N = dif * dif.T
58     self.cov = weight_old * self.cov + weight_N * (cov_intra_N + cov_inter_N)
59
60 def cov_and_mean_N(self, seqs_N):
61     '''
62     Calculate the mean and intra-group covariance of the sequences to be added to
63     the data-set.
64     '''
65     seqs_T = seqs_N.T
66     mean_N = np.zeros(self.Lcad_B)
67     for i in range(self.Lcad_B):
68         mean_N[i] = np.mean(seqs_T[i])
69         cov_N = np.cov(seqs_T, bias = True)
70     return(mean_N, cov_N)
71
72 @nb.njit(parallel = False, fastmath = True)
73 def calc_weight(n_seq_i, N):
74     '''
75     Calculate the weights based on the number of sequences considered and the number
76     of sequences to be added in this step.
77     '''
78     old_Nseq = n_seq_i - N
79     old_weight = old_Nseq / n_seq_i
80     N_weight = N / n_seq_i
81     return(old_weight, N_weight)

```

6.5 Sequence binary and int8 transformation.

In this annex is presented a Python script that contains all the functions needed to transform the data base format and the sequences structure.

```

1 import numpy as np
2 import numba as nb
3 from Bio.Seq import Seq
4 from Bio import SeqIO
5 from Bio.SeqRecord import SeqRecord
6
7
8 aminosdict = ['A', 'C', 'D', 'E', 'F', 'G', 'H', 'I', 'K', 'L', 'M', 'N', 'P', 'Q',
9              'R', 'S', 'T', 'V', 'W', 'Y', '-']
10 aminosdictnum = [65, 67, 68, 69, 70, 71, 72, 73, 75, 76, 77, 78, 80, 81,
11                 82, 83, 84, 86, 87, 89, 45]
12 aminosdict_np = np.asarray(aminosdict)
13 aminosdictnum_np = np.asarray(aminosdictnum)
14
15 Namin = len(aminosdict) - 1
16
17 @nb.njit(parallel = True, fastmath = True)
18 def floattoint_vect(Cadenas, Limite):
19     '''
20     Going from numbers between 0 and 1 to 0 or 1 in an array of sequences.

```

```

21     '''
22     ncad, lcad = Cadenas.shape
23     for j in nb.prange(ncad):
24         for i in nb.prange(lcad):
25             if (Cadenas[j,i] > Limite):
26                 Cadenas[j,i] = 1
27             else:
28                 Cadenas[j,i] = 0
29     return Cadenas
30
31 @nb.njit(parallel = False, fastmath = True)
32 def floattoint(Cadena, Limite):
33     '''
34     Going from numbers between 0 and 1 to 0 or 1 in a single sequence.
35     '''
36     lcad = len(Cadena)
37     for i in range(lcad):
38         if (Cadena[i] > Limite):
39             Cadena[i] = 1
40         else:
41             Cadena[i] = 0
42     return Cadena
43
44
45 def escribe_fasta(Chains, Ident, file):
46     '''
47     From a list of sequences and another list with the identifiers, generate a
48     fasta file.
49     '''
50     for N_C in range(len(Chains)):
51         record = SeqRecord(Seq(str(Chains[N_C])), id=Ident[N_C],
52                             description="Semilla")
53         SeqIO.write(record, file, 'fasta')
54
55 def ref_seqs_mat(RefSeqsList):
56     '''
57     Creates an array of reference sequences from a list of reference
58     sequences generated by biopython.
59     '''
60     nrefseqs = len(RefSeqsList)
61     lrefseqs = len(RefSeqsList[0].seq)
62     Ref_Seqs = np.chararray((nrefseqs, lrefseqs), unicode = True)
63     for nseq, ref in enumerate(RefSeqsList):
64         Ref_Seqs[nseq] = ref.seq
65     return(Ref_Seqs)
66
67 @nb.njit(parallel = False, fastmath = True)
68 def ascii_to_binary(seq):
69     '''
70     Receives a single sequence in ascii and returns that sequence
71     translated to binary
72     '''
73     bseq = np.zeros(len(seq) * Namin, dtype = np.int8)
74     for pos, amin in enumerate(seq):
75         if amin != '-':
76             for i in range(Namin):
77                 if amin == aminsdict_np[i]:
78                     aminc = i
79                     bseq[pos*Namin + aminc] = 1

```

```

80         break
81     return(bseq)
82
83 @nb.njit(parallel = False, fastmath = True)
84 def ascii_to_binary_vect(seq):
85     '''
86     Receives a matrix sequence in ascii and returns that matrix
87     translated to binary
88     '''
89     dim = seq.shape
90     bseq = np.zeros((dim[0], dim[1] * Namin), dtype = np.int8)
91     for s in range(dim[0]):
92         for pos, amin in enumerate(seq[s]):
93             if amin != '-':
94                 for i in range(Namin):
95                     if amin == aminosdict_np[i]:
96                         aminc = i
97                         bseq[s, pos*Namin + aminc] = 1
98                         break
99     return(bseq)
100
101 @nb.njit(parallel = False, fastmath = True)
102 def uint8_to_binary(seq):
103     '''
104     Recibe una unica secuencia en uint8 y devuelve esa secuencia
105     traducidas a binario
106     '''
107     bseq = np.zeros(len(seq) * Namin, dtype = np.int8)
108     for pos, num in enumerate(seq):
109         if num != 45:
110             for i in range(Namin):
111                 if num == aminosdictnum_np[i]:
112                     aminc = i
113                     bseq[pos*Namin + aminc] = 1
114                     break
115     return(bseq)
116
117 @nb.njit(parallel = False, fastmath = True)
118 def uint8_to_binary_vect(seq):
119     '''
120     Receives an array in which each row is a sequence in
121     uint8 and returns those sequences translated to binary
122     '''
123     dim = seq.shape
124     bseq = np.zeros((dim[0], dim[1] * Namin), dtype = np.int8)
125     for s in range(dim[0]):
126         for pos, num in enumerate(seq[s]):
127             if num != 45:
128                 for i in range(lbloque + 1):
129                     if num == aminosdictnum_np[i]:
130                         aminc = i
131                         bseq[s, pos*lbloque + aminc] = 1
132                         break
133     return(bseq)
134
135
136 def ascii_to_unit8_vect(seqs):
137     '''
138     Recibe una matriz en la que cada fila es una secuencia en

```

```

139     ascii y devuelve esas secuencias traducidas a uint8
140     '''
141     dim = seqs.shape
142     mat = np.zeros((dim[0],dim[1]))
143     for i in range(dim[0]):
144         mat[i] = np.asarray(seqs[i], dtype = '|S1').view(np.uint8).astype(int)
145     return(mat)
146
147 def ascii_to_unit8(seq):
148     '''
149     Recibe una secuencia en ascii y devuelve esa secuencia
150     traducida a uint8
151     '''
152     mat = np.asarray(seq, dtype = '|S1').view(np.uint8).astype(int)
153     return(mat)
154
155
156 @nb.njit(parallel = False, fastmath = True)
157 def binary_to_uint8_vect(bseq, Lcad):
158     '''
159     Recibe una matriz en la que cada fila es una secuencia en
160     binario y devuelve esas secuencias traducidas a uint8
161     '''
162     dim = bseq.shape
163     seqs_uint = np.zeros((dim[0],Lcad), dtype = np.uint8)
164     for s in range(dim[0]):
165         for bloque in range(Lcad):
166             one = -1
167             for i in range(bloque*Namin, ((bloque + 1)*Namin)):
168                 if bseq[s, i] == 1:
169                     one = i - bloque*Namin
170                     break
171             if one == -1:
172                 seqs_uint[s, bloque] = int(aminosdictnum_np[20])
173             else:
174                 seqs_uint[s, bloque] = int(aminosdictnum_np[one])
175     return(seqs_uint)
176
177 @nb.njit(parallel = False, fastmath = True)
178 def binary_to_uint8(bseq):
179     '''
180     Recibe una secuencia en binario y devuelve esa secuencia
181     traducida a uint8
182     '''
183     Lcad = int(len(bseq)/Namin)
184     seq_uint = np.zeros(Lcad, dtype = np.uint8)
185     for bloque in range(Lcad):
186         one = -1
187         for i in range(bloque*Namin, ((bloque + 1)*Namin)):
188             if bseq[i] == 1:
189                 one = i - bloque*Namin
190                 break
191         if one == -1:
192             seq_uint[bloque] = int(aminosdictnum_np[20])
193         else:
194             seq_uint[bloque] = int(aminosdictnum_np[one])
195     return(seq_uint)
196
197 def binary_to_ascii(seq_b):

```

```

198     '''
199     Recibe una secuencia en binario y devuelve esa secuencia
200     traducida a ascii
201     '''
202     lcad_b = len(seq_b)
203     lcad = int(lcad_b/Namin)
204     seq = np.chararray(lcad, unicode = True)
205     for bloque in range(lcad):
206         ind = np.where(seq_b[bloque*Namin:((bloque + 1)*Namin)] == 1)[0]
207         if ind.size > 0:
208             seq[bloque] = aminosdict[int(ind)]
209         else:
210             seq[bloque] = aminosdict[Namin]
211     return seq
212
213 def binary_to_ascii_vect(seq_b):
214     '''
215     Recibe una matriz en la que cada fila es una secuencia en
216     binario y devuelve esas secuencias traducidas a ascii
217     '''
218     nseq, lcad_b = seq_b.shape
219     lcad = int(lcad_b/Namin)
220     seq = np.chararray((nseq, lcad), unicode = True)
221     for seq_i in range(nseq):
222         for bloque in range(lcad):
223             ind = np.where(seq_b[seq_i, bloque*Namin:((bloque + 1)*Namin)] == 1)[0]
224             if ind.size > 0:
225                 seq[seq_i, bloque] = aminosdict[int(ind)]
226             else:
227                 seq[seq_i, bloque] = aminosdict[Namin]
228     return seq
229
230 def unit8_to_ascii(seq):
231     '''
232     Recibe una unica secuencia en uint8 y devuelve esa secuencia
233     traducidas a ascii
234     '''
235     lcad = len(seq)
236     mat = np.chararray(lcad, unicode = True)
237     for i in range(lcad):
238         mat[i] = chr(seq[i])
239     return(mat)
240
241 def unit8_to_ascii_vect(seq):
242     '''
243     Recibe una matriz en la que cada fila es una secuencia en
244     ascii y devuelve esas secuencias traducidas a uint8
245     '''
246     nseq, lcad = seq.shape
247     mat = np.chararray((nseq, lcad), unicode = True)
248     for j in range(nseq):
249         for i in range(lcad):
250             mat[j, i] = chr(seq[j, i])
251     return(mat)

```

6.6 d_{ant} , d_{inv} and covariance norm calculation.

In this annex is presented a Python script that calculates d_{ant} , d_{inv} and covariance norm calculation while adding sequences to the data set as is described in Section 2.5.

```
1 import numpy as np
2 import numba as nb
3 from Bio import SeqIO
4 import matplotlib.pyplot as plt
5 import sys
6
7 sys.path.append("/home/david/Desktop/TFM/seqsTFM-base/seqsTFM/")
8 from large_cov import *
9 from transform import ref_seqs_mat, ascii_to_binary_vect, floattoint_vect
10 from stats import prior,
11
12 ##### Functions #####
13
14 def carga_seq(n_seq_i):
15     '''
16     Add sequences from the sequences list.
17     '''
18     if Igor:
19         if n_seq_i == Nseq_min:
20             N_seqs = ascii_to_binary_vect(ref_seqs_mat(seqs_list[:n_seq_i]))
21         else:
22             N_seqs = ascii_to_binary_vect(ref_seqs_mat(seqs_list[(n_seq_i - N):
23                 n_seq_i]))
24     if Julia_C:
25         if n_seq_i == Nseq_min:
26             N_seqs = Main.gen_delt_sem(n_seq_i, Lcad_B, gauss_obj)
27         else:
28             N_seqs = Main.gen_delt_sem( N , Lcad_B, gauss_obj)
29         if Discreto:
30             N_seqs = floattoint_vect( N_seqs , Limite)
31
32     return( N_seqs )
33
34 def cov_ref_fija(data, cov_i, ind):
35     '''
36     Chech if the covariance matrix is invertible.
37     '''
38     try:
39         np.linalg.inv(cov_i)
40     except:
41         return None, None, True
42     else:
43         return cov_i, cov_i, False
44
45 @nb.njit(parallel = False, fastmath = True)
46 def calcula_dists(cov_i, cov_ref, cov_ant):
47     '''
48     Calculates d1_ant, d2_ant, d1_inv, d2_inv and the covariance norm.
49     '''
50     norm_ant = np.linalg.norm(cov_ant)
51     d_ant = np.linalg.norm(cov_i - cov_ant)/norm_ant
52     norm_ref = np.linalg.norm(cov_ref)
53     d_ref = np.linalg.norm(cov_i - cov_ref)/norm_ref
54     d_t = np.linalg.det((cov_i + cov_i.T)/2)
```

```

54 d_i = np.linalg.norm(cov_i)
55 return d_ref, d_ant, d_i, d_t
56
57 def itera( , old , oldU):
58     data = np.zeros((5, N_pasos + 1))
59     cov_ref = None
60     check = True
61     cov_t = System(Lcad_B = Lcad_B)
62
63     for n_seq_i in range(Nseq_min, Nseq + N , N):
64
65         print('Trabajando con: ', n_seq_i, 'seqs.', end = '\r')
66         ind = int((n_seq_i-Nseq_min)/ N)
67         data[0, ind] = n_seq_i
68
69         N_seqs = carga_seq(n_seq_i)
70         cov_t.update( N_seqs , n_seq_i, N , old , oldU, )
71         if check:
72             cov_ref, cov_ant, check = cov_ref_fija(data, cov_t.cov, ind)
73         else:
74             data[1, ind], data[2, ind], data[3, ind], data[4, ind] = calcula_dists(cov_t.
75             cov, cov_ref, cov_ant)
76             cov_ant = cov_t.cov
77         return data
78 ##### Parametros #####
79
80 Igor = False
81 if not Igor:
82     Julia_C = True
83 else:
84     Julia_C = False
85
86 Discreto = True
87 Limite = 0.367
88
89 N = 10000
90 Nseq_min = 10000
91
92     = 0.01
93
94 ##### Main code #####
95 if Julia_C:
96     Lcad_B = 2980
97     Lcad = 149
98     Nseq = 1300000
99     print('Max seq: ', Nseq)
100
101 from julia.api import Julia
102 jl = Julia(compiled_modules = False)
103 from julia import Main
104
105 jl = Main.include('../Opt_lambda/GenSem.jl')
106
107 ref = np.loadtxt('../DataBases/Modelo/Ref/MedBin.txt')
108 ref = ref[:2980]
109 ref = np.loadtxt('../DataBases/Modelo/Ref/CovBin.txt')
110 ref = ref[:2980, :2980]
111 print( ref.shape, ref.shape)

```

```

112     ref_norm = np.linalg.norm( ref )
113     ref_norm = np.linalg.norm( ref )
114     print('Modelo cargado')
115
116     gauss_obj = Main.gen_obj_gauss( ref , ref )
117     print('Obj. gauss generado')
118
119 if Igor:
120     seqs_list = list(SeqIO.parse(open('../DataBases/Igor/1M_354k.fasta'), 'fasta'))
121     Lcad = len(seqs_list[0].seq)
122     Lcad_B = Lcad*20
123     Nseq = len(seqs_list)
124     print('Max seq: ', Nseq)
125
126 old , oldU = prior(Lcad, Lcad_B)
127 print('Prior calculada')
128
129 N_pasos = int((Nseq - Nseq_min)/ N) + 1
130 print('Num. pasos: ', N_pasos)
131
132 data = itera( , old , oldU)
133 h = data.shape[1] - 1
134
135 plt.plot(data[0, 1:h], data[1, 1:h], label = 'Norm_ref')
136 plt.plot(data[0, 1:h], data[2, 1:h], label = 'Norm_ant')
137 plt.xlabel('Nseq')
138 plt.legend()
139 plt.show()
140 plt.plot(data[0, 1:h], data[3, 1:h])
141 plt.plot(data[0, 1:h], data[4, 1:h])
142
143 with open('Data.txt', 'w') as file:
144     np.savetxt(file, data)
145
146 plt.xlabel('Nseq')
147 plt.ylabel('Cov Norm')
148 plt.show()

```

6.7 Nucleotide sequence extraction

In this annex is presented two Python scripts. The first one is able to extract the nucleotide sequences from the data obtained from [1]. The second one reproduce the Figure 7 of [1] and divide the sequences into three clusters. This three clusters are the high density zones presented in Section 3.3.

```

1 from Bio.Seq import Seq
2 from Bio import SeqIO
3 from Bio.SeqRecord import SeqRecord
4 import numpy as np
5
6
7 def describe_fasta(Chains, Ident, file):
8     '''
9     From a list of sequences and another list with the identifiers, generate a
10     fasta file.
11     '''
12     for N_C in range(len(Chains)):

```

```

13     record = SeqRecord(Seq(Chains[N_C]), id = Ident[N_C], description = "Asti")
14     SeqIO.write(record, file, 'fasta')
15
16 def cabecera(line_sp, ident, ref_dict):
17     '''
18     Comprueba si la linea es una cabecera del tipo 'Query=....' y si
19     esta contenida en la referencia.
20     '''
21     if line_sp[0] == 'Query=':
22         ident = line_sp[1]
23         if ident.split(':')[0] in ref_dict:
24             ident = line_sp[1]
25         else:
26             ident = None
27         return ident
28     else: return ident
29
30 def lee_file(file_name, ref_dict):
31     with open(file_name, 'r') as file:
32         check = False
33         ident = None
34         cab = 0
35         seq = ''
36         IGHV1_2_ident = []
37         seq_list = []
38         ident_list = []
39         for line in file:
40             line_sp = line.split()
41             if len(line_sp) > 0:
42                 ident = cabecera(line_sp, ident, ref_dict)
43                 if line_sp[0].split('*')[0] == 'IGHV1-2' and line_sp[2].split('*')[0] == '
IGHJ2' and line_sp[6] == 'Yes':
44                     if ident != None:
45                         check = True
46
47                 if check:
48                     if len(line_sp) > 4:
49                         if line_sp[0] == 'Total':
50                             # print(line_sp)
51                             IGHV1_2_ident.append(line_sp[7])
52                         if line_sp[0].split('_')[0] == 'Query':
53                             seq += line_sp[2]
54
55                         if line_sp[0] == 'Effective':
56                             check = False
57                             seq_list.append(seq)
58                             ident_list.append(ident)
59                             seq = ''
60
61         file.close()
62         print('secuencias = ', len(seq_list))
63         print('ident = ', len(ident_list))
64         print(len(IGHV1_2_ident))
65     return seq_list, ident_list, IGHV1_2_ident
66
67 def cambia_cab(ini_list):
68     '''
69     Modifica las cabeceras para que no aparezca la multiplicidad.
70     '''

```

```

71 cab_new = []
72 for item in ini_list:
73     cab = item.id.split(':')[0]
74     cab_new.append(cab)
75 return cab_new
76
77 def filtro_buenas(ref_list):
78     '''
79     Secuencias del archivo de Asti buenas.
80     '''
81     i = 0
82     for item in ref_list:
83         if item.id.split('.')[0] == 'SRR277211':
84             i += 1
85     print('Filtradas:',i)
86     return i
87
88 def comprueba_mult(ref_list):
89     '''
90     Checkea cuantas secuencias tienen multiplicidad 1 y cuanto
91     suma la multiplicidad total del data base.
92     '''
93     i = 0
94     j = 0
95     ref_cab_mult = []
96     ref_cab_no_mult = []
97     for item in ref_list:
98         item_s = item.id.split(':')
99         if item.id.split('.')[0] == 'SRR277211':
100             i += int(item_s[1])
101             ref_cab_mult.append(item.id)
102             ref_cab_no_mult.append(item_s[0])
103             if int(item_s[1]) == 1:
104                 j += 1
105     print('Con multiplicidad 1:', j)
106     print('Mult. total:',i)
107     return ref_cab_mult, ref_cab_no_mult
108
109 def sustituye_cab(ref_cab, ig_cab, no_mult_cab):
110     cab = []
111     for item in ig_cab:
112         temp = item.split(':')[0]
113         if temp in no_mult_cab:
114             for ide in ref_cab:
115                 if ide.split(':')[0] == temp:
116                     cab.append(ide)
117     return cab
118
119 def genera_base_not_unique(ident_list, seq_list, IGHV1_2_ident):
120     ident_list_fin = []
121     seq_list_fin = []
122     for j,item in enumerate(ident_list):
123         rep = int(item.split(':')[1])
124         for i in range(rep):
125             ident_list_fin.append(ident_list[j] + '_' + IGHV1_2_ident[j])
126             seq_list_fin.append(seq_list[j])
127     print('n_cad_final = ',len(ident_list_fin))
128     return ident_list_fin, seq_list_fin
129

```

```

130 def a_ade_eq(ident_list, eq_list):
131     for i in range(len(ident_list)):
132         ident_list[i] += '_' + eq_list[i]
133     return ident_list
134
135 #####
136 ##### Main Code #####
137 #####
138
139 ref_file = '../DatosAstatic/ModAsti/allYES_uniqueAA_c90-140_IGHV1-2
         _IGHJ2_SEQMA45_3SE9GG_reH2_allSEQMA0.afasta'
140 ref_dict = SeqIO.to_dict(SeqIO.parse(open(ref_file), 'fasta'))
141 print('File Asti:', len(ref_dict))
142 ref_list = list(SeqIO.parse(open(ref_file), 'fasta'))
143
144 new_id = cambia_cab(ref_list)
145
146 buenas = filtro_buenas(ref_list)
147
148 ref_cab_mult, ref_cab_no_mult = comprueba_mult(ref_list)
149
150 file_name = '../DatosAstatic/DataBruscolini3/unique_INVCOMPL_NOBAR_all.igblast'
151 # file_name = 'Igbblast_cor.txt'
152 seq_list, ident_list, IGHV1_2_ident = lee_file(file_name, new_id)
153 print(len(ident_list))
154 print(len(IGHV1_2_ident))
155
156
157 ident_list_f = a_ade_eq(ident_list, IGHV1_2_ident)
158 file_out = 'Unique_Fig_7.fasta'
159 with open(file_out, 'w') as file:
160     escribe_fasta(seq_list, ident_list_f, file)
161
162 final_cab = sustituye_cab(ref_cab_mult, ident_list, ref_cab_no_mult)
163 print(len(final_cab))
164 ident_list_mult, seq_list_fin = genera_base_not_unique(final_cab, seq_list,
        IGHV1_2_ident)
165 file_out = 'Not_Unique_Fig_7.fasta'
166 with open(file_out, 'w') as file:
167     escribe_fasta(seq_list_fin, ident_list_mult, file)

```

```

1 from Bio import SeqIO
2 import numpy as np
3 import matplotlib.pyplot as plt
4 import matplotlib.cm as cm
5 import seaborn as sn
6 from Bio.SeqRecord import SeqRecord
7 from Bio.Seq import Seq
8
9 class Sequence:
10
11     def __init__(self, idt, seq):
12         self.id = idt
13         self.seq = seq
14         self.IGHV1_2 = 0
15         self.VRC_PG04 = 0
16
17     def add_sim_IGHV1_2(self, sim_IGHV1_2):
18         self.IGHV1_2 = sim_IGHV1_2

```

```

19
20 def add_sim_VRC_PG04(self, sim_VRC_PG04):
21     self.VRC_PG04 = sim_VRC_PG04
22
23 def minus_to_mayus(ref_seq):
24     mayus = ''
25     for i in ref_seq:
26         if i == 'a': mayus += 'A'
27         elif i == 'c': mayus += 'C'
28         elif i == 'g': mayus += 'G'
29         elif i == 't': mayus += 'T'
30         else: print('Caracter extra o:', i)
31     return mayus
32
33 def extrae_V_sim(seq_list):
34     sim = np.zeros(len(seq_list))
35     for j, item in enumerate(seq_list):
36         sim[j] = float((item.id.split('_')[1]).split('%')[0])/100
37     return sim
38
39 def list_data(seq_list, IGHV1_2_ident, VRC_PG04_ident):
40     items = []
41     for i in range(len(seq_list)):
42         item = Sequence((seq_list[i].id).split('_')[0], seq_list[i].seq)
43         item.add_sim_IGHV1_2(IGHV1_2_ident[i])
44         item.add_sim_VRC_PG04(float(VRC_PG04_ident[i]))
45         items.append(item)
46     return items
47
48 def Heat_map(IGHV1_2, VRC_PG04):
49     """
50     Calcula el mapa de calor.
51     """
52     mat = np.zeros((2, len(IGHV1_2)))
53     mat[0] = IGHV1_2
54     mat[1] = np.asarray(VRC_PG04)
55     sn.heatmap(mat, cmap='Blues')
56
57 def ah_shit_here_we_go_again(file_name):
58     ident_list = []
59     ident = None
60     nuc_ident = []
61     sim = np.zeros(2)
62     with open(file_name, 'r') as file:
63         for line in file:
64             line_sp = line.split()
65             if len(line_sp) > 0:
66
67                 if line_sp[0] == 'Query=':
68                     ident_list.append(line_sp[1])
69
70                 if len(line_sp) > 6 and line_sp[0] == 'Identities':
71                     if ident == ident_list[-1]:
72                         div = line_sp[2].split('/')
73                         sim[0] += int(div[0])
74                         sim[1] += int(div[1])
75                     else:
76                         div = line_sp[2].split('/')
77                         sim[0] = int(div[0])

```

```

78         sim[1] = int(div[1])
79         ident = ident_list[-1]
80
81         if line_sp[0] == 'Effective':
82             nuc_ident.append(sim[0]/sim[1])
83
84         # print(len(ident_list))
85         # print(len(nuc_ident))
86         return ident_list, nuc_ident
87
88 def Plotea_3D(IGHV1_2_ident, VRC_PG04_ident):
89
90     fig = plt.figure()          #create a canvas, tell matplotlib it's 3d
91     ax = fig.add_subplot(111, projection='3d')
92
93     hist, xedges, yedges = np.histogram2d(IGHV1_2_ident, VRC_PG04_ident, bins=(20,20)
94         , range = [[0.6, 1], [0.6, 1]])
95     xpos, ypos = np.meshgrid(xedges[:-1]+xedges[1:], yedges[:-1]+yedges[1:])
96
97     xpos = xpos.flatten()/2.
98     ypos = ypos.flatten()/2.
99     zpos = np.zeros_like (xpos)
100
101     dx = xedges [1] - xedges [0]
102     dy = yedges [1] - yedges [0]
103     dz = hist.flatten()
104
105     cmap = cm.get_cmap('jet') # Get desired colormap - you can change this!
106     max_height = np.max(dz)   # get range of colorbars so we can normalize
107     min_height = np.min(dz)
108     # scale each z to [0,1], and get their rgb values
109     rgba = [cmap((k-min_height)/max_height) for k in dz]
110
111     ax.bar3d(xpos, ypos, zpos, dx, dy, dz, color=rgba, zsort='average')
112     ax.invert_xaxis()
113     plt.xlabel("Nucleotides identity with IGHV1-2*02")
114     plt.ylabel("Nucleotides identity with VRC_PG04")
115     plt.savefig("Fig_7_Asti.pdf")
116     plt.show()
117
118 def decide_clust_nuc(item):
119     dist_hiper = np.sqrt((0.75 - item.IGHV1_2)**2 + (0.95 - item.VRC_PG04)**2)
120     dist_germ = np.sqrt((0.95 - item.IGHV1_2)**2 + (0.725 - item.VRC_PG04)**2)
121     dist_inter = np.sqrt((0.75 - item.IGHV1_2)**2 + (0.75 - item.VRC_PG04)**2)
122     return ('hiper', 'germ', 'inter')[np.argmin((dist_hiper, dist_germ, dist_inter))]
123
124 def div_clust_nuc(data):
125     hiper = []
126     germ = []
127     inter = []
128     print('len data:', len(data))
129     for ind, item in enumerate(data):
130         temp_ident = item.id
131         clust = decide_clust_nuc(item)
132         if clust == 'hiper':
133             if not hiper or temp_ident != hiper[-1].id:
134                 hiper.append(item)
135         elif clust == 'germ':
136             if not germ or temp_ident != germ[-1].id:

```

```

136     germ.append(item)
137     elif clust == 'inter':
138         if not inter or temp_ident != inter[-1].id:
139             inter.append(item)
140     print(len(hiper))
141     print(len(germ))
142     print(len(inter))
143     print()
144     return hiper, germ, inter
145
146
147 def extrae_data(items, ref_data, ref_list_cab):
148     ident_clus = []
149     ident_tot = []
150     final_items = []
151
152     for item in items:
153         ident_clus.append(item.id)
154     for item in ref_list_cab:
155         ident_tot.append(item)
156
157     for id_c in ident_clus:
158         ind = ident_tot.index(id_c.split(':')[0])
159         final_items.append(ref_data[ind])
160     return final_items
161
162 def escribe_fasta(data, file):
163     Chains = []
164     Ident = []
165     for i in range(len(data)):
166         Chains.append(data[i].seq)
167         Ident.append(data[i].id)
168     print(len(Chains))
169     for N_C in range(len(Chains)):
170         record = SeqRecord(Seq(str(Chains[N_C])), id = Ident[N_C])
171         SeqIO.write(record, file, 'fasta')
172 def quita_ident(ref_data):
173     ident = []
174     for item in ref_data:
175         ident.append(item.id.split('_')[0].split(':')[0])
176     return ident
177
178 ref_file = './filtradas_all_yes.fasta'
179 ref_list = list(SeqIO.parse(open(ref_file), 'fasta'))
180 ref_list_cab = quita_ident(ref_list)
181
182 seq_file = './Unique_Fin.fasta'
183 seq_list = list(SeqIO.parse(open(seq_file), 'fasta'))
184
185 # ref_mayus = minus_to_mayus(ref_list[1].seq)
186 # print(ref_mayus)
187 # exit()
188 file_name = 'VRC_PG04_align.txt'
189 ident_list, VRC_PG04_ident = ah_shit_here_we_go_again(file_name)
190
191 file_name = 'IGHV1_2_02_align.txt'
192 ident_list, IGHV1_2_ident = ah_shit_here_we_go_again(file_name)
193
194 IGHV1_2_ident_s = extrae_V_sim(seq_list)

```

```

195
196 IGHV1_2_ident = np.asarray(IGHV1_2_ident)
197 VRC_PG04_ident = np.asarray(VRC_PG04_ident)
198
199 data = list_data(seq_list, IGHV1_2_ident, VRC_PG04_ident)
200
201 Plotea_3D(IGHV1_2_ident, VRC_PG04_ident)
202 # exit()
203
204 hiper, germ, inter = div_clust_nuc(data)
205 print(len(hiper + germ + inter))
206 hiper_prot = extrae_data(hiper, ref_list, ref_list_cab)
207 print(len(hiper_prot))
208
209 s = 0
210 i = 0
211 for item in hiper:
212     s += item.IGHV1_2
213     i += item.VRC_PG04
214 print(s/len(hiper))
215 print(i/len(hiper))
216
217 germ_prot = extrae_data(germ, ref_list, ref_list_cab)
218 print(len(germ_prot))
219
220 s = 0
221 i = 0
222 for item in germ:
223     s += item.IGHV1_2
224     i += item.VRC_PG04
225 print(s/len(germ))
226 print(i/len(germ))
227
228 inter_prot = extrae_data(inter, ref_list, ref_list_cab)
229 print(len(inter_prot))
230
231 s = 0
232 i = 0
233 for item in inter:
234     s += item.IGHV1_2
235     i += item.VRC_PG04
236 print(s/len(inter))
237 print(i/len(inter))
238
239
240 file_hiper = 'hiper_mut_clust.fasta'
241 # escribe_fasta(hiper_prot, file_hiper)
242 with open('hiper_mut_clust.fasta', 'w') as file:
243     SeqIO.write(hiper_prot, file, 'fasta')
244
245 file_germ = 'germ_clust.fasta'
246 # escribe_fasta(germ_prot, file_germ)
247 with open('germ_clust.fasta', 'w') as file:
248     SeqIO.write(germ_prot, file, 'fasta')
249
250 file_inter = 'inter_clust.fasta'
251 # escribe_fasta(inter_prot, file_inter)
252 with open('inter_clust.fasta', 'w') as file:
253     SeqIO.write(inter_prot, file, 'fasta')

```

```

254
255 # plt.hist(a, bins = 100, color = 'red')
256 # plt.show()
257 # sim = extrae_V_sim(seq_list)
258 # plt.hist(IGHV1_2_ident, bins = 100, color = 'red')
259 # plt.xlabel('Nucleotides identity with IGHV1-2*02')
260 # plt.show()
261 # plt.hist(VRC_PG04_ident, bins = 100, color = 'red')
262 # plt.xlabel('Nucleotides identity with VRC_PG04')
263 # plt.show()

```

6.8 Kullback-Leibler divergence calculation.

In this annex is presented a Python script that calculates the Kullback-Leibler divergence from one data set to another reference data set.

```

1 from Bio import SeqIO
2 import numpy as np
3 import numba as nb
4 import sys
5 import matplotlib.pyplot as plt
6
7 sys.path.append("/home/david/Desktop/TFM/seqsTFM-base/seqsTFM/")
8 from transform import ascii_to_unit8_vect, ref_seqs_mat
9
10 aminosdictnum = [65, 67, 68, 69, 70, 71, 72, 73, 75, 76, 77, 78, 80, 81,
11                  82, 83, 84, 86, 87, 89, 45]
12
13 amin_dict_u8 = np.asarray(aminosdictnum)
14
15 class_input = [('seqs_mat', nb.float64[:, :]), ('n_seq', nb.int64), ('l_seq', nb.
16             int64)]
17 @nb.experimental.jitclass(class_input)
18 class seqs_mat:
19     def __init__(self, seqs_mat):
20         self.seqs_mat = seqs_mat
21         self.n_seq = seqs_mat.shape[0]
22         self.l_seq = seqs_mat.shape[1]
23
24 @nb.njit(parallel = False, fastmath = True)
25 def rep_amin(col_vect, amin_ref, n_seq):
26     n_rep = 0
27     for amin in col_vect:
28         if amin == amin_ref:
29             n_rep += 1
30     return n_rep/n_seq
31
32 @nb.njit(parallel = True, fastmath = True)
33 def probabilities(seqs_mat_u8):
34     seqs_mat_u8_t = seqs_mat_u8.seqs_mat.T
35     prob_mat = np.zeros((len(amin_dict_u8), seqs_mat_u8.l_seq))
36     for col in range(seqs_mat_u8.l_seq):
37         for ind, amin in enumerate(amin_dict_u8):
38             prob_mat[ind, col] = rep_amin(seqs_mat_u8_t[col], amin, seqs_mat_u8.n_seq)
39     return prob_mat
40
41 @nb.njit(parallel = False, fastmath = True)

```

```

42 def curva_KL(q, p, l_seq, fallo):
43     KL = np.zeros(l_seq)
44     h = 0 #Cuenta las columnas en las que q_i(aa)=0 pero p_i(aa) no se anula.
45     rep_null = np.zeros(l_seq)
46     prob_null = np.zeros(l_seq)
47     for pos in range(l_seq):
48         for amin in range(len(amin_dict_u8)):
49             if q[amin, pos] == 0:
50                 if p[amin, pos] == 0:
51                     KL[pos] += 0
52                 else:
53                     h += 1
54                     rep_null[pos] += 1
55                     KL[pos] += p[amin, pos]*np.log(p[amin, pos]/fallo) #si q_i(aa)=0 lo pongo
56                                     #germline, para evitar que explote.
57                     prob_null[pos] += p[amin, pos]
58                     if p[amin, pos] > 0.25:
59                         print('null')
60                         print(p[amin, pos])
61                         print(pos, amin_dict_u8[amin])
62                         print()
63                 else:
64                     if p[amin, pos] == 0:
65                         KL[pos] += 0
66                     else:
67                         KL[pos] += p[amin, pos]*np.log(p[amin, pos]/q[amin, pos])
68     print('fallos:', h)
69     return KL, rep_null, prob_null
70
71 germ_list = list(SeqIO.parse('../HammingDist/Plots_def/Comp/C2_1.fasta', 'fasta'))
72 # print(len(germ_list))
73 germ_seqs_mat = seqs_mat(ascii_to_unit8_vect(ref_seqs_mat(germ_list)))
74 germ_prob = probabilities(germ_seqs_mat)
75
76 inter_list = list(SeqIO.parse('../HammingDist/Plots_def/Comp/C2_2.fasta', 'fasta'))
77 # print(len(inter_list))
78 inter_seqs_mat = seqs_mat(ascii_to_unit8_vect(ref_seqs_mat(inter_list)))
79 inter_prob = probabilities(inter_seqs_mat)
80
81 hiper_list = list(SeqIO.parse('../HammingDist/Plots_def/Comp/C1.fasta', 'fasta'))
82 # print(len(hiper_list))
83 hiper_seqs_mat = seqs_mat(ascii_to_unit8_vect(ref_seqs_mat(hiper_list)))
84 hiper_prob = probabilities(hiper_seqs_mat)
85
86 print('Intermedio:')
87 KL_1, rep_null_1, prob_null_1 = curva_KL(q = germ_prob, p = inter_prob, l_seq =
    germ_seqs_mat.l_seq, fallo = 1/len(germ_list))
88 print('10 mejores posiciones inter:')
89 print(np.argsort(KL_1*-1)[:10])
90
91 print('Hipermutado:')
92 KL_2, rep_null_2, prob_null_2 = curva_KL(q = germ_prob, p = hiper_prob, l_seq =
    germ_seqs_mat.l_seq, fallo = 1/len(germ_list))
93 print('10 mejores posiciones hiper:')
94 print(np.argsort(KL_2*-1)[:10])
95
96 x = range(len(KL_1))
97

```

```

98 figure_config()
99 plt.plot(x, KL_1, label = 'C2_2 cluster')
100 plt.plot(x, KL_2, label = 'C1 cluster')
101 plt.legend(loc = 'upper right')
102 plt.xlabel('Sequence position')
103 plt.ylabel('Kullback Leibler divergence')
104 plt.savefig('KL_curve.pdf')
105 plt.show()
106
107 figure_config()
108 plt.vlines(x=63, ymin=0, ymax=8, color = 'black', label = 'Seq. pos. = 63')
109 plt.vlines(x=73, ymin=0, ymax=8, color = 'purple', label = 'Seq. pos. = 73')
110 plt.plot(x[50:80], KL_1[50:80], label = 'C2_2 cluster')
111 plt.plot(x[50:80], KL_2[50:80], label = 'C1 cluster')
112 plt.legend(loc = 'upper left')
113 plt.xlabel('Sequence position')
114 plt.ylabel('Kullback Leibler divergence')
115 plt.savefig('KL_curve_amp.pdf')
116 plt.show()

```

6.9 Minimize α

In this annex is presented a Python script that was used to calculate the relaxation matrix of the OU process, as it is explained in the Section 3.4.

```

1 import numba as nb
2 import numpy as np
3 import sys
4 from math import pi as
5 from math import lgamma
6 from scipy.optimize import minimize
7 from scipy.linalg import expm
8 from Bio import SeqIO
9
10 sys.path.append("../seqsTFM-base/seqsTFM")
11 from transform import ref_seqs_mat, ascii_to_binary_vect
12 from stats import prior,
13 from large_cov import *
14 from plot import plot
15
16 ##### Classes #####
17
18 class Alt_System(System):
19     '''
20     Heredamos la clase System y a adimos el numero de secuencias
21     '''
22     def __init__(self, Lcad_B, n_seqs):
23         System.__init__(self, Lcad_B)
24         self.n_seqs = n_seqs
25
26 ##### Functions #####
27
28 def E_diag(alpha, t):
29     D = -alpha*t
30     for l in range(Lcad):
31         D[l*Q:(l + 1)*Q, l*Q:(l + 1)*Q] = np.exp(D[l*Q:(l + 1)*Q, l*Q:(l + 1)*Q])
32     return D
33

```

```

34 def ineq_constraint(x):
35     """constrain all elements of x to be >= 0"""
36     return x
37
38 def eq_constraint(x):
39     """constrain the sum of all rows to be equal to 1"""
40     return x - 1
41
42 def from_m_to_v_cajas_sim(m, L, Q):
43     triang = int(Q*(Q + 1)/2)
44     v = np.zeros(L*triang, dtype = np.float64)
45     for i in range(L):
46         caja = m[i*Q:(i+1)*Q, i*Q:(i+1)*Q]
47         v[triang*i:triang*(i+1)] = caja[np.triu_indices(Q)]
48     return v
49
50 def from_v_to_m_cajas_sim(v, L, Q):
51     triang = int(Q*(Q + 1)/2)
52     m = np.zeros((L*Q, L*Q))
53     caja = np.zeros((Q, Q))
54     for i in range(L):
55         caja[np.triu_indices(Q)] = v[i*triang:(i + 1)*triang]
56         caja[np.tril_indices(Q, -1)] = caja.T[np.tril_indices(Q, -1)]
57         m[i*Q:(i + 1)*Q, i*Q:(i + 1)*Q] = caja
58     return m
59
60 def carga_seq(n_seq_i, seqs_list, N, Nseq_min):
61     '''
62     Carga parte de las secuencias, en lugar de todas, de la lista para evitar
63     overflows.
64     '''
65     if n_seq_i == Nseq_min:
66         N_seqs = ascii_to_binary_vect(ref_seqs_mat(seqs_list[:n_seq_i]))
67     elif n_seq_i > len(seqs_list):
68         N_seqs = ascii_to_binary_vect(ref_seqs_mat(seqs_list[(n_seq_i - N):]))
69     else:
70         N_seqs = ascii_to_binary_vect(ref_seqs_mat(seqs_list[(n_seq_i - N):n_seq_i
71         ])))
72     return (N_seqs)
73
74 def cov_and_mean_with_prior(seqs, n_seq, old, oldU, N, Nseq_min):
75     '''
76     Calcula la media y la covarianza de una lista de secuencias, aadiendo una prior
77     .
78     '''
79     cov_and_mean = Alt_System(Lcad_B, n_seq)
80     for n_seq_i in range(Nseq_min, n_seq + N, N):
81         N_seqs = carga_seq(n_seq_i, seqs, N, Nseq_min)
82         cov_and_mean.update(N_seqs, n_seq_i, N, old, oldU, )
83     return cov_and_mean
84
85 def cov_and_mean_no_prior(seqs, n_seq, N, Nseq_min):
86     '''
87     Calcula la media y la covarianza de una lista de secuencias.
88     '''
89     cov_and_mean = Alt_System(Lcad_B, n_seq)
90     for n_seq_i in range(Nseq_min, n_seq + N, N):
91         N_seqs = carga_seq(n_seq_i, seqs, N, Nseq_min)
92         cov_and_mean.update_no_prior(N_seqs, n_seq_i, N)

```

```

90     return cov_and_mean
91
92 def omega_theta_0(IGoR_system, Germ_system):
93     init_point = Alt_System(Lcad_B, IGoR_system.n_seqs)
94     _0 = IGoR_system.n_seqs
95     n_0 = Germ_system.n_seqs
96     init_point.mean_t = (n_0/(n_0 + _0))*Germ_system.mean_t + (_0/(n_0 + _0))*
        IGoR_system.mean_t
97
98     d_mean = (Germ_system.mean_t - IGoR_system.mean_t).reshape(Lcad_B, 1)
99     init_point.cov = (IGoR_system.cov + n_0*Germ_system.cov + ((n_0*_0)/(n_0 + _0
        ))*(d_mean*d_mean.T))/(_0 + G + 2)
100
101     return init_point
102
103 def extrae_ _ t(var_opt):
104     result = var_opt.x
105     t = result[:2]
106     _v = result[2:]
107     = from_v_to_m_cajas_sim(_v, Lcad, Q)
108     return t,
109
110 def save_ _ t(t, , ):
111     name_t = 't_' + str(int( )) + '.txt'
112     with open(name_t, 'w') as file:
113         np.savetxt(file, t)
114
115     name_t = ' _ ' + str(int( )) + '.txt'
116     with open(name_t, 'w') as file:
117         np.savetxt(file, )
118
119 def scores_klus_k( , , nseq_k):
120     _prima = + nseq_k
121     _prima = + nseq_k
122
123     ln_num = (G/2)*np.log(_prima) + (G*2)*(_prima + G + 2)*np.log(_prima + G +
        2) - (G/2)*(_prima + G + 2)
124     ln_den = (G/2)*np.log(2* ) + (G*_prima/2)*np.log(2) + lgamma(_prima/2) -
        lgamma(G)
125     return(ln_num - ln_den)
126
127 def prob_ac(t, , , System_1, System_2, System_3, pre_med, pre_cov, _ref,
    _ref):
128     s_1 = scores_klus_k( , , System_1.n_seqs)
129     s_2 = scores_klus_k( , , System_2.n_seqs)
130     s_3 = scores_klus_k( , , System_3.n_seqs)
131
132     F = (G + 2)/2
133     print(s_1 + s_2 + s_3)
134
135     ldet_1 = log_det_ _func( , t[0], pre_cov, pre_med, , System_1, _ref, _ref
        )
136     ldet_2 = log_det_ _func( , t[1], pre_cov, pre_med, , System_2, _ref, _ref
        )
137     ldet_3 = log_det_ _func( , 1, pre_cov, pre_med, , System_3, _ref, _ref)
138
139     ldet_t = F*(ldet_1 + ldet_2 + ldet_3)
140     print(ldet_t)
141     return (ldet_t, s_1 + s_1 + s_1 - ldet_t)

```

```

142
143 def log_det_ _func( , t , pre_cov, pre_med, , System_k, _ref , _ref):
144
145     E = expm(- * t)
146     = _ref - np.matmul(E, np.matmul(pre_cov, E.T))
147
148     pre_med = (pre_med).reshape(Lcad_B, 1)
149     = ( _ref - np.matmul(E, pre_med).reshape(Lcad_B,))
150
151     inc = (System_k.mean_t - ).reshape(Lcad_B, 1)
152     _prima = + System_k.n_seqs*System_k.cov + ((System_k.n_seqs* )/(System_k.
        n_seqs + ))*(np.matmul(inc, inc.T))
153
154     return (np.linalg.slogdet( _prima)[1])
155
156 def min_ (variables, , _ref , _ref , pre_cov, pre_med, k1_System, k2_System,
        k3_System, germ_System):
157     '''
158     Calcula el score siguiendo la ecuacion 34 y los maximiza buscando maximizar la
        ecuacion 30
159     '''
160     t_1 = variables[0]
161     t_2 = variables[1]
162
163     B = from_v_to_m_cajas_sim(variables[2:], Lcad, Q)
164     = np.matmul(B.T, B)
165
166
167     F = (G + 2)/2
168
169     ldet_1 = log_det_ _func( , t_1, pre_cov, pre_med, , k1_System, _ref , _ref
        )
170     ldet_2 = log_det_ _func( , t_2, pre_cov, pre_med, , k2_System, _ref , _ref
        )
171     ldet_3 = log_det_ _func( , 1, pre_cov, pre_med, , k3_System, _ref , _ref)
172     ldet_4 = log_det_ _func( , 0, pre_cov, pre_med, , germ_System, _ref , _ref
        )
173
174     # ldet_t = F*(ldet_1 + ldet_2 + ldet_3)
175     ldet_t = (ldet_1 + ldet_2 + ldet_3 + ldet_4)
176     return ldet_t
177
178 def itera_ ( _ref , _ref , Init_point, k1_System, k2_System, k3_System,
        germ_System):
179     '''
180     Calcula el alpha optimo para cada valor de lambda.
181     '''
182     variables = np.ones(int((Q*(Q + 1)/2)*Lcad + 2))*0.1
183     # variables = np.ones(200 + 2)
184
185     pre_med = _ref - Init_point.mean_t
186     pre_cov = _ref - Init_point.cov*(G + 2 + Init_point.n_seqs)
187
188     variables[2:] = from_m_to_v_cajas_sim(np.copy( _ref), Lcad, Q) #Inicializamos
        alpha plano.
189
190     bounds_1 = ((1e-10, None), (1e-10, None))
191     bounds_2 = ((None, None), (None, None))*int((Q*(Q + 1)/2)*Lcad/2)
192     # bounds_2 = ((None, None), (None, None))*int(100)

```

```

193     bounds = bounds_1 + bounds_2
194
195     # cons = ({'type': 'eq', 'fun': lambda x: x[5] - 1})
196
197     _v = np.asarray([1123])
198
199     for _ in _v:
200
201         = + G - 1
202
203         log_det_ini, score_ini = prob_ac(np.ones(3), np.ones((200,200)), , ,
k1_System, k2_System, k3_System, pre_med, pre_cov, _ref, _ref)
204         print('score_ini =', score_ini, 'log_det_ini = ', log_det_ini)
205
206         var_op = minimize(min_ , (variables), args = ( , _ref, _ref, pre_cov,
pre_med, k1_System, k2_System, k3_System, germ_System), bounds = bounds,
207             options = {"disp":True, 'maxfun': 35000, 'maxiter': 5})
208
209         print(var_op.success)
210         print(var_op.x[:2])
211         t_op, B_op = extrae_ _ t(var_op)
212         _op = np.matmul(B_op.T, B_op)
213         save_ _ t(t_op, _op, )
214         log_det_fin, score_fin = prob_ac(t_op, _op, , , k1_System, k2_System,
k3_System, pre_med, pre_cov, _ref, _ref)
215         print('times:', t_op)
216         print('score_fin =', score_fin)
217         print('log_det_fin =', log_det_fin)
218
219     ##### Parameters #####
220     G = 200
221     Q = 20
222
223     Lcad = 10
224     Lcad_B = Lcad*Q
225
226     _prior = 0.01
227
228     N_Igor = 10000
229     Nseq_min_Igor = 10000
230
231     N_Asti = 100
232     Nseq_min_Asti = 100
233
234     ##### Main Code #####
235
236     ## Carga de datos ##
237
238     #Datos Igor#
239     Igor_list = list(SeqIO.parse(open('../HammingDist/Plots_def/outkluster_75_cort.
fasta'), 'fasta'))
240
241     #Germline (cluster CC2_1)#
242     germ_list = list(SeqIO.parse(open('../HammingDist/Plots_def/Cortadas/CC2_1.fasta'),
'fasta'))
243
244     #Resto de clusters (CC1, CC2_2_1, CC2_2_2)#
245     k1_list = list(SeqIO.parse(open('../HammingDist/Plots_def/Cortadas/CC2_2_1.fasta'),
'fasta'))

```

```

246 k2_list = list(SeqIO.parse(open('../HammingDist/Plots_def/Cortadas/CC2_2_2.fasta'),
    'fasta'))
247 k3_list = list(SeqIO.parse(open('../HammingDist/Plots_def/Cortadas/CC1.fasta'),
    'fasta'))
248
249 print('Datos cargados.')
250 #Prior plana para calcular la cov de IGoR#
251 old , oldU = prior(Lcad, Lcad_B)
252 print('Prior calculada')
253
254 #Calculo de matrices de covarianza de IGoR y germline#
255 Igor_System = cov_and_mean_with_prior(Igor_list, len(Igor_list), old , oldU,
    _prior, N_Igor, Nseq_min_Igor)
256 Germ_System = cov_and_mean_no_prior(germ_list, len(germ_list), N_Asti ,
    Nseq_min_Asti)
257
258 #Calculo de omega_0 y theta_0#
259 Init_point = omega_theta_0(Igor_System, Germ_System)
260
261 #Calculo de matrices de covarianza resto de clusters (CC1, CC2_2_1, CC2_2_2)#
262 k1_System = cov_and_mean_no_prior(k1_list, len(k1_list), N_Asti , Nseq_min_Asti)
263 k2_System = cov_and_mean_no_prior(k2_list, len(k2_list), N_Asti , Nseq_min_Asti)
264 k3_System = cov_and_mean_no_prior(k3_list, len(k3_list), N_Asti , Nseq_min_Asti)
265
266 print('Covarianzas halladas.')
267 itera_(oldU, old , Init_point, k1_System, k2_System, k3_System, Germ_System)
268 print('Scores optimos hallados.')

```