

Introducción a la Bioinformática

Búsquedas de secuencias en bases de datos

Heurísticas, Hits, significancia de los resultados

Fernán Agüero

fernan@iib.unsam.edu.ar

Instituto de Investigaciones Biotecnológicas

Universidad Nacional de San Martín



Búsqueda de secuencias por similitud

- Tenemos un **método (algoritmo)** que nos garantiza un alineamiento óptimo entre dos secuencias
- Tenemos un **sistema de scoring** complejo que refleja mejor nuestras ideas biológicas acerca de lo que es un alineamiento
- Tenemos una **secuencia de interés** (*query*)
- Y una **base de datos** con secuencias
- Cómo implementaríamos una búsqueda **por similitud** contra una base de datos de secuencias?

Método de fuerza bruta = Búsqueda exhaustiva

- Buscar secuencias similares al query en la base de datos
- Tomar **una por una** las secuencias de la base de datos
- Calcular un **alineamiento y su score**
- **Ordenar los alineamientos** en base al score
- Finalmente usar **nuestro criterio** y evaluar si las secuencias al tope del ranking son lo suficientemente similares

https://en.wikipedia.org/wiki/Brute-force_search

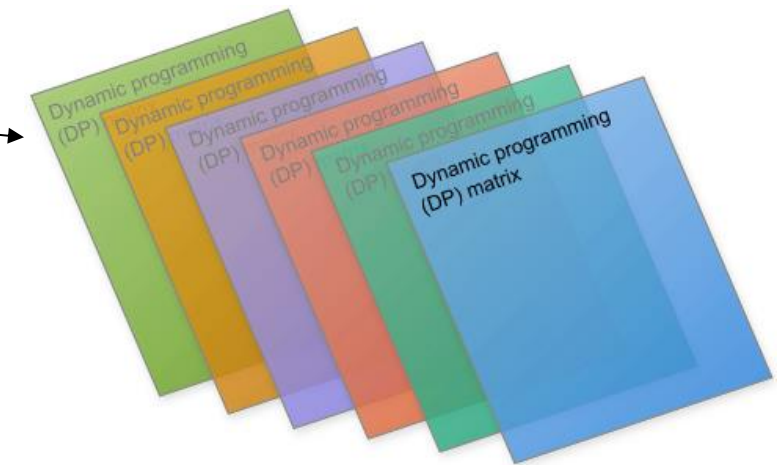
Espacio de búsqueda

- Cuántos **espacios de búsqueda** reconocibles hay?
- Las búsquedas de secuencias por similitud son “exactas” si recorren completamente el espacio de búsqueda



Espacios de búsqueda

- Hay dos espacios de búsqueda reconocibles:
 - El espacio de todas **las secuencias de la base de datos**
 - Database size (number of sequences)
 - El espacio de todos **los alineamientos posibles entre dos secuencias**
 - Dynamic programming matrix ($M \times N$)
 - M = longitud de la secuencia query
 - N = longitud de la secuencia n de la base de datos
- Las búsquedas de secuencias por similitud son “exactas” si recorren en forma exhaustiva ambos espacios
 - Ej: Smith-Waterman sobre toda la base de datos



- **A continuación vamos a introducir heurísticas para reducir estos espacios de búsqueda**
 - **Estrategias heurísticas para filtrar la base de datos**
 - **Heurísticas para reducir el espacio de alineamientos posibles que se explora efectivamente**

heurística

La **heurística** es un conjunto de estrategias, reglas prácticas o atajos mentales que permiten resolver problemas, tomar decisiones rápidas y encontrar soluciones de manera eficiente cuando no se dispone de tiempo o información completa. Su origen etimológico proviene del griego *heuriskein*, que significa «hallar» o «inventar».  Diccionario de la lengü... +3

<https://es.wikipedia.org/wiki/Heurística>

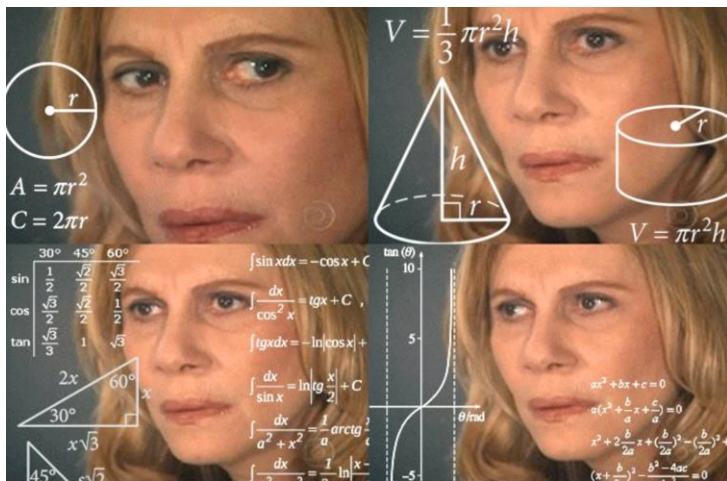
Heurísticas para reducir el espacio de secuencias

MLIKRDELVISWASHERE

query sequence (ejemplo)

Vamos a intentar reducir el espacio de **las secuencias de la base de datos**
Database size (number of sequences)

Ideas? Heuriskein?



Heurísticas para reducir el espacio de secuencias

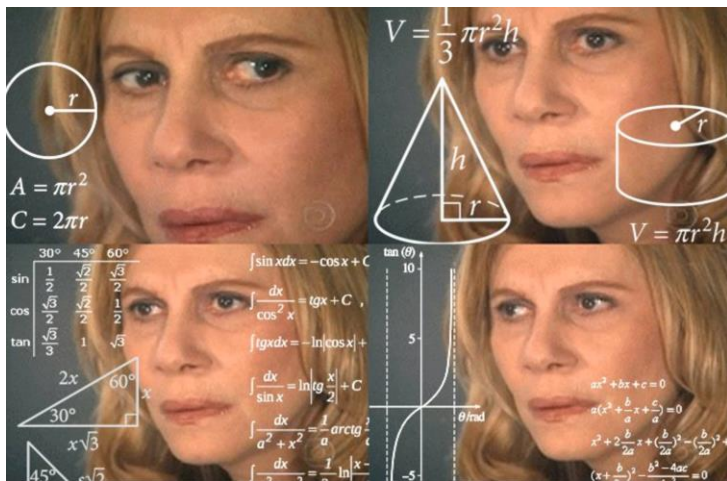
MLIKRDELVISWASHERE

query sequence (ejemplo)

Vamos a intentar reducir el espacio de **las secuencias de la base de datos**

Database size (number of sequences)

Ideas? Heuriskein?



Ideas:

- Descartar secuencias (saltarlas)
- Cuáles?
 - Las que no empiezan con Met?
 - Las que tienen una composición de aminoácidos diferente?
 - Las que son mucho mas largas?
 - Mucho mas cortas?

Busquemos inspiración en **los libros**.

Los índices ayudan a encontrar cosas.

Que es un indice?

- Un mapa que conecta una **palabra** con una **página**
- Lookup tables, Hashes
 - Conectan datos



Hashes ~ Indices: ejemplo con anotaciones de genes

- Ejemplo de índices en bases de datos: buscar todos los registros que contengan la palabra **‘kinase’** en la descripción de la secuencia

gi	acc	def
214734077	20	Xenopus laevis rhodopsin
123456789	56789	Mus musculus casein kinase

Ejemplo mínimo.

Tamaño real de la base de datos:
Cientos de millones de registros (10⁸)

- Indexar la columna ‘def’ genera una nueva tabla

word	list of GIs
casein	1234, 3245, 43678, 123456 ...
kinase	432, 5678, 32456, 123456 ...
laevis	36314, 214734, ...
mus	23467, 98732, 123456, 312456, 567
muscu	123467, 98732, 123456, 567983 ...
rhodopsin	214734, 223466, 873212, 23587, 29
xenopus	28462, 36314, 98476, 214734 ...

Tamaño del índice?
Cuántos registros tiene?

Hashing methods

Hashing es un método común para acelerar búsquedas en bases de datos.

Compilar un “diccionario” o ‘lookup table’ de palabras a partir de la secuencia ‘query’. Arma un índice con todas las posibles sub-palabras contenidas en la secuencia.

Longitud de la secuencia: 19

Cantidad de palabras: 17

MLIIKRDELVISWASHERE

query sequence

MLI
LII
IIK
IKR
KRD
RDE
DEL
ELV
LVI
VIS
ISW
SWA
WAS
ASH
SHE
HER
ERE

todas las palabras
posibles de longitud 3

ktup = 3
Word size = 3

Secuencia Query

Seq: “MLIIKRDELVISWASHERE”

Longitud: 19

Ktup: 1; Cantidad de palabras: 19

Ktup: 2; Cantidad de palabras : 18

Ktup: 3; Cantidad de palabras : 17

Secuencias en la Base de Datos

Seq: “MLIIKRDELVISWASHERE”

Longitud: 19

Nro de Secuencias posibles: $20^{19} \sim 5.24 \cdot 10^{24}$

Ktup 1; Cantidad de palabras?

Ktup: 2; Cantidad de palabras?

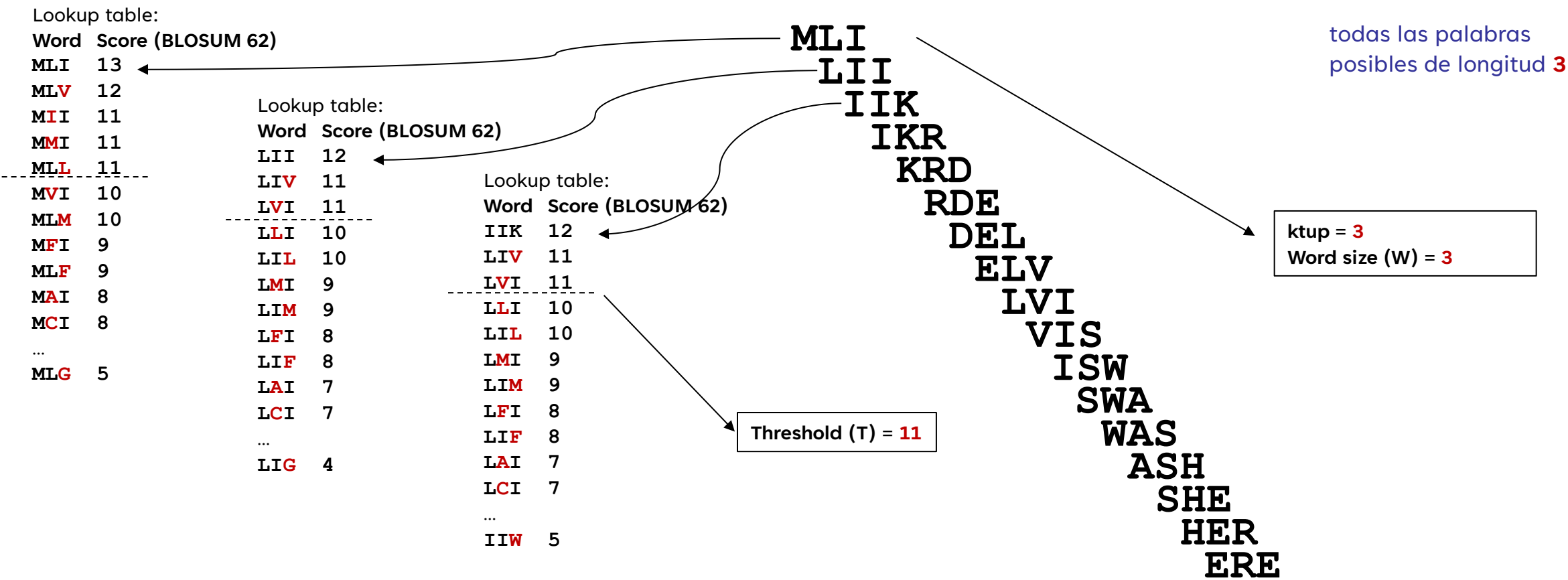
Ktup: 3; Cantidad de palabras?

Ejemplo de lookup table de secuencias

Word	{seq,pos}	Acc	Def	Prot Seq
AAA	{L07770,232}; ...	L0770	Xenopus laevis rhodopsin mRNA, complete cds	MNGTEGPNFYVPM SNKTGVVRSPFDYPQY YLAEPWQYSALAA YMFL LILLGLPINFMTLF VTIQHKKLRTPLNYILLNLVFANHF MVLCG FTVTMYTSMHGYFIFGQTGCYIEGFFATLG GEVALWSLVVLAVERYMVVCKPMANFRFG ENHAIMGVAFTWIMALSCAAPPLFGWSRYI PEGMQCSCGV DYYTLKPEVNNESFVIYMF I VHFTIPLIVIFFCYGRLLCTVKEAA AQQQES ATTQKAEKEVTRMVVIMVVF LICWVPYAY VAFYIFTHQGSNFGPVFMTVPAFFAKSSAIY NPVIYIVLNKQFRNCLITTLCCGKNPFGDED GSSAATSKTEASSVSSSQVSPA
AAC				
AAD				
AAE				
AAF				
AAG				
AAH				
AAI				
AAK				
AAL	{L0771,221}; ...			
AAM		L0771	Xenopus laevis (clone pXga1) transducin alpha subunit mRNA, complete cds	MGAGASAE EKHSRELEKKLKEDADKDART VKLLLLGAGESGKSTIVKQM KIIHQDGYSVE ECLEFIAIIYGNTLQSMIAIVKAMNTLNIQFG DPA RQDDARKLMHLADTIDEGSMPKEMSDI IGRLWKDTGIQACFDRASEYQLNDSAGYYL NDLDR LVIPGYVPTEQDVLRSRVKTTGIIET QFGLKDLNFRMFDVGGQRSERKKWIHC FE GVTCIIFI AAL SAYDMVLVEDDEVNRMHESL HLFNSICNHRYFATTSIVLFLNKKDVFTEKIK KAHLSICFPDYDGPNTYEDAGNYIKTQFLE LNMRRDVKEIYSHMTCATDTENVKFV FDAV TDIIKENLKD CGLF
AAN				
AAP	{L0770; 168}; ...			
AAQ	{BC002171,77}; ...			
AAR				
AAS	{BC002171,81}; ...			
AAT	{L0770; 336}			
AAV				
AAW				
AAZ				
AAZ		BC002171	Mus musculus casein kinase 1, alpha 1, mRNA (cDNA clone IMAGE:3486713), partial cds.	KKQKYEKISEKKMSTPVEVLCKGFPAEFAM YLN YCRGLRFEEAPDYMYLRQLFRILFRTL NHQYDYTFDWTMLKQKAA QQAASSSGQG QQAQTPTGKQTDKTKSNMKG F
AAZ				
AAZ		...		
...				

Busquedas usando lookup tables

Lookup table: todas las posibles sub-palabras de longitud **Ktup** que coincidan con un **score** $\geq T$



Secuencia Query

Seq: “**MLIIKRDELVISWASHERE**”

Longitud: **19**

Ktup: **3**; Cantidad de palabras : **17** →

Secuencias en la Base de Datos

RefSeq (Julio 2026) = 482,864,455 (proteínas)

Parámetros:
Ktup = **3**
Threshold = **11**

Ktup: **3**; Cantidad de palabras = max 8000

Ktup: **3**; Secuencias con palabras con score $\geq$ **11**?

Cuántas secuencias con **17** palabras con score $\geq$ 11?

Cuántas secuencias con **16** palabras con score $\geq$ 11?

Cuántas secuencias con **15** palabras con score $\geq$ 11?

Cuántas secuencias con **14** palabras con score $\geq$ 11?

Cuántas secuencias con **13** palabras con score $\geq$ 11?

Cuántas secuencias con **12** palabras con score $\geq$ 11?

...

Cuál es la frecuencia de coincidencias de k-meros en pares de secuencias biológicas conservadas?

Table 3. Library search sensitivity^a

PROTEINS			Related – missed/unrelated – found			
Query	Superfamily	Family size	BLAST	ktup = 2	ktup = 1 opt	Smith-Waterman
HAHU	Hemoglobin	505	35	43 –	17 +	16 +
K1HUAG	Ig κ V region	280	54	100 –	17 +	21 +
OOHU	G-protein CR	165	27	42 –	26 +	22 +
CCHU	Cytochrome c	142	25	29 –	27 –	27 –
N2KF1U	Snake neurotoxin	109	3	5 –	4 –	4 –
XURT8C	Glutathione transferase	106	8	20 –	8	8
TPHUCS	Calcium binding	106	8	17 –	7 +	11 –
OKHU2C	Protein kinase	97	54	52 +	37 +	37 +
FEPE	Ferredoxin	93	53	54 –	53	53
RKMDS	Ribulose-bisphosphate carboxylase	77	0	1 –	0	0
K3HU	Ig κ C region	74	22	26 –	18 +	16 +
HMIVV	Hemagglutinin	73	1	4 –	0 +	0 +
HLHUB2	Histocompatibility Ag	71	2	15 –	1 +	0 +
IPHU	Insulin	69	3	3	3	3
CYBOA	α-Crystallin	67	7	8 –	4 +	4 +
PSHU	Phospholipase A2	58	2	2	2	2
DEHUGL	G3-PDH	46	0	1 –	1 –	1 –
TVHURA	Transforming protein (N-ras)	45	1	1	1	1

Pearson WR. Comparison of methods for searching protein sequence databases. Protein Sci. 1995 Jun;4(6):1145-60. doi: 10.1002/pro.5560040613. PMID: 7549879; PMCID: PMC2143149.

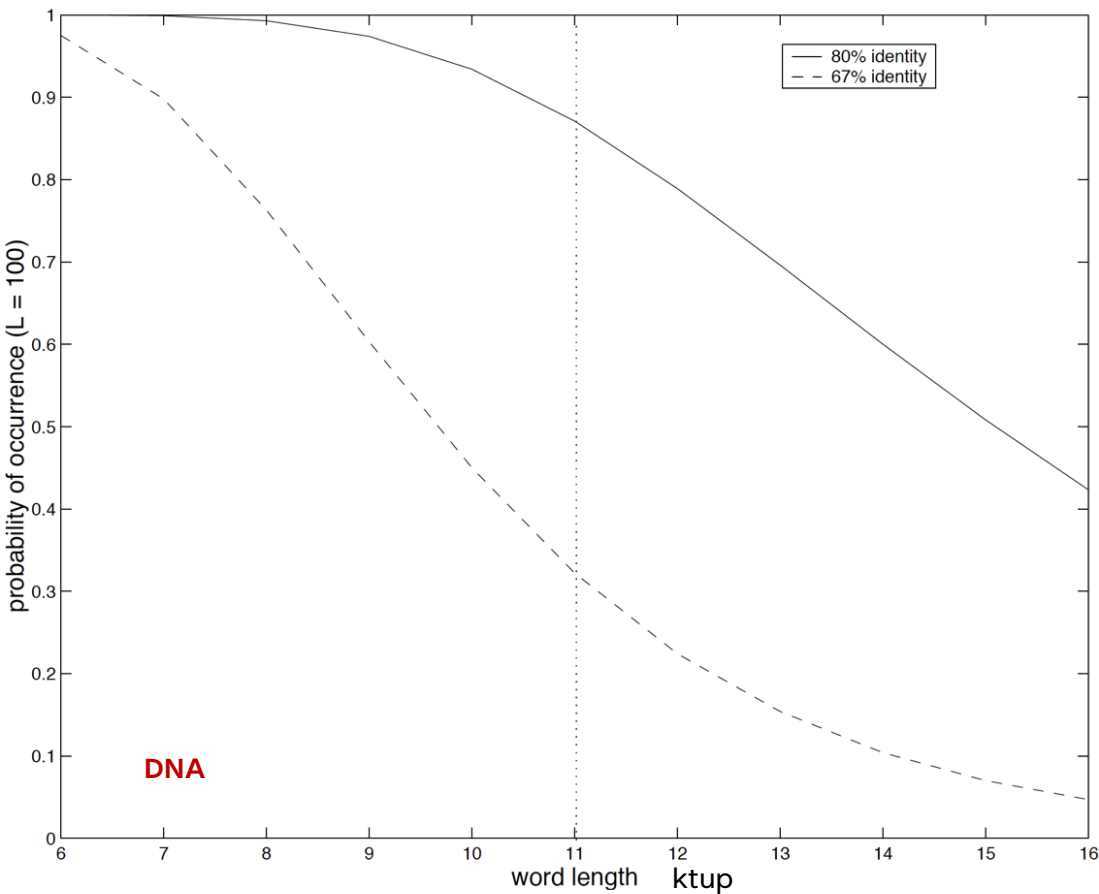
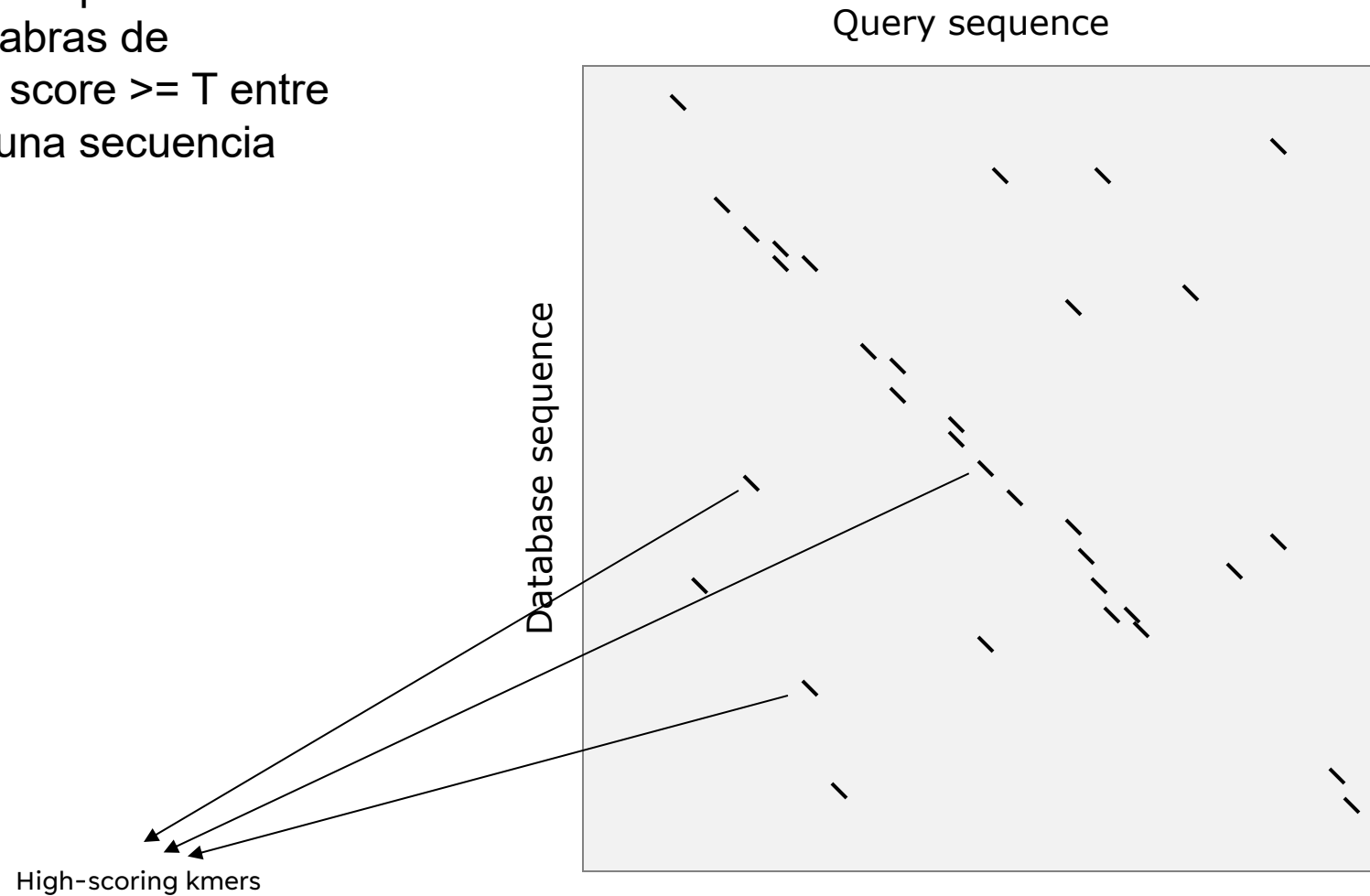


Gráfico de longitud de kmeros (eje x) vs probabilidad de ocurrencia para secuencias de ADN de longitud = 100 (eje y) y alineamientos de 80% o 67% de identidad.

https://gepstaging.wustl.edu/lessons/Smith_Waterman_to_BLAST/Smith_Waterman_to_BLAST.html

High-Scoring Words Map

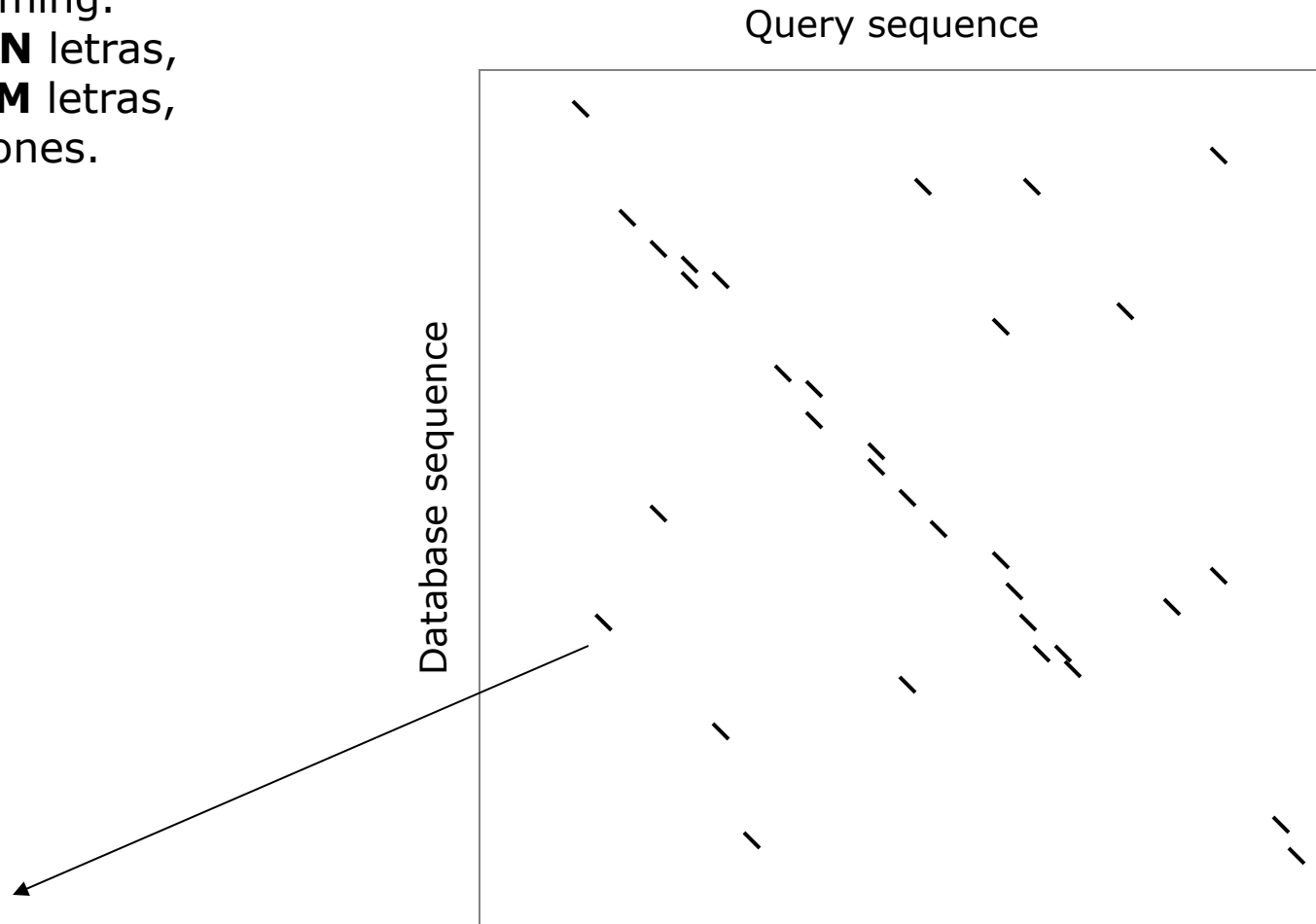
El mapa de palabras compartidas muestra todas las palabras de longitud = k tup, y con $\text{score} \geq T$ entre la secuencia query y una secuencia de la base de datos.



High-Scoring Words Map

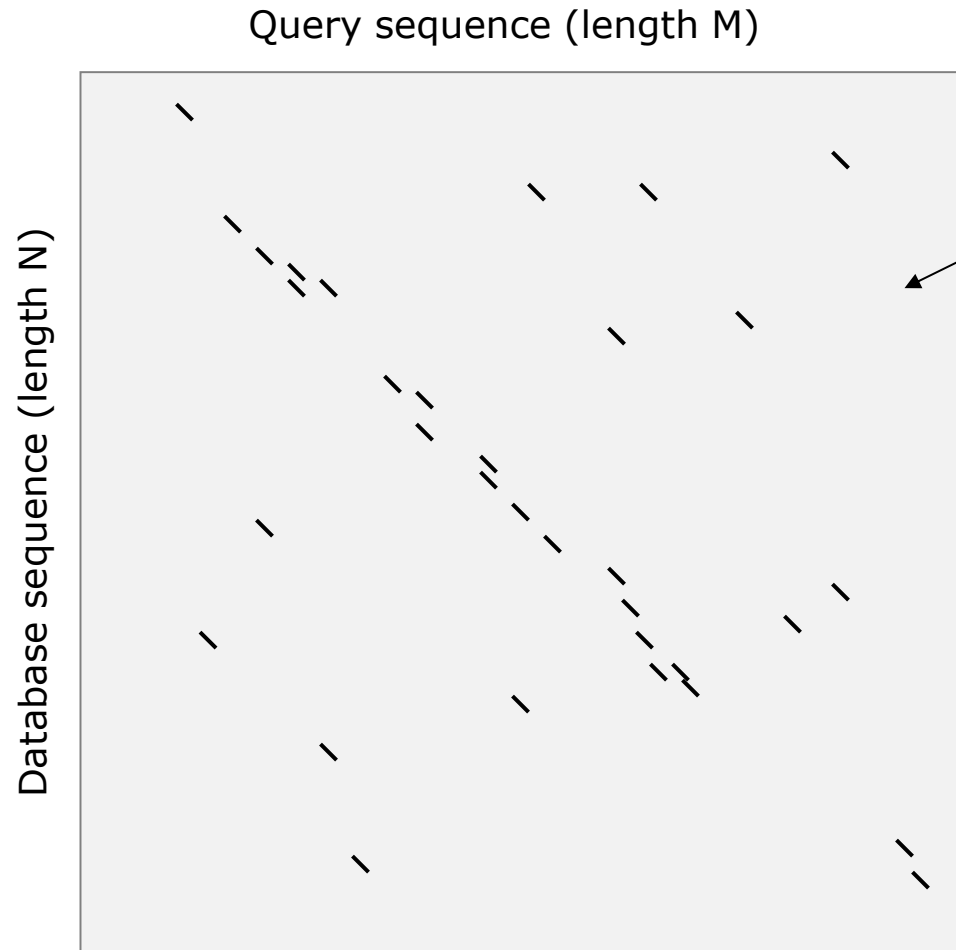
La búsqueda más simple es un gran ejemplo de dynamic programming.
Para una secuencia query de **N** letras,
contra una base de datos de **M** letras,
se requieren **MxN** comparaciones.

Cómo reducir este espacio de búsqueda?



High-scoring kmer

Heurísticas para reducir espacio de búsqueda

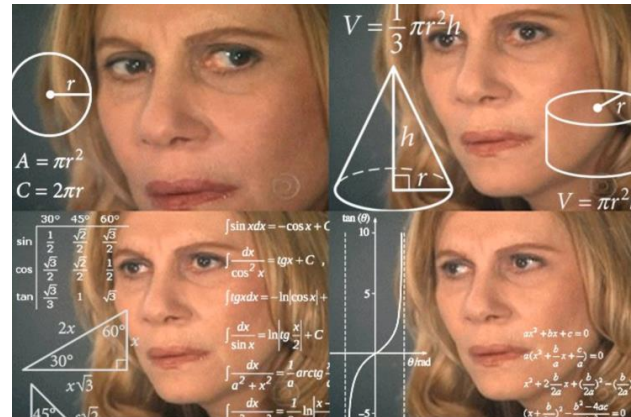


Ahora vamos a intentar reducir

El espacio de todos los **alineamientos posibles entre dos secuencias**

- **Dynamic programming matrix (M x N)**
 - **M** = longitud de la secuencia query
 - **N** = longitud de la secuencia de la base de datos

Cómo? Ideas?

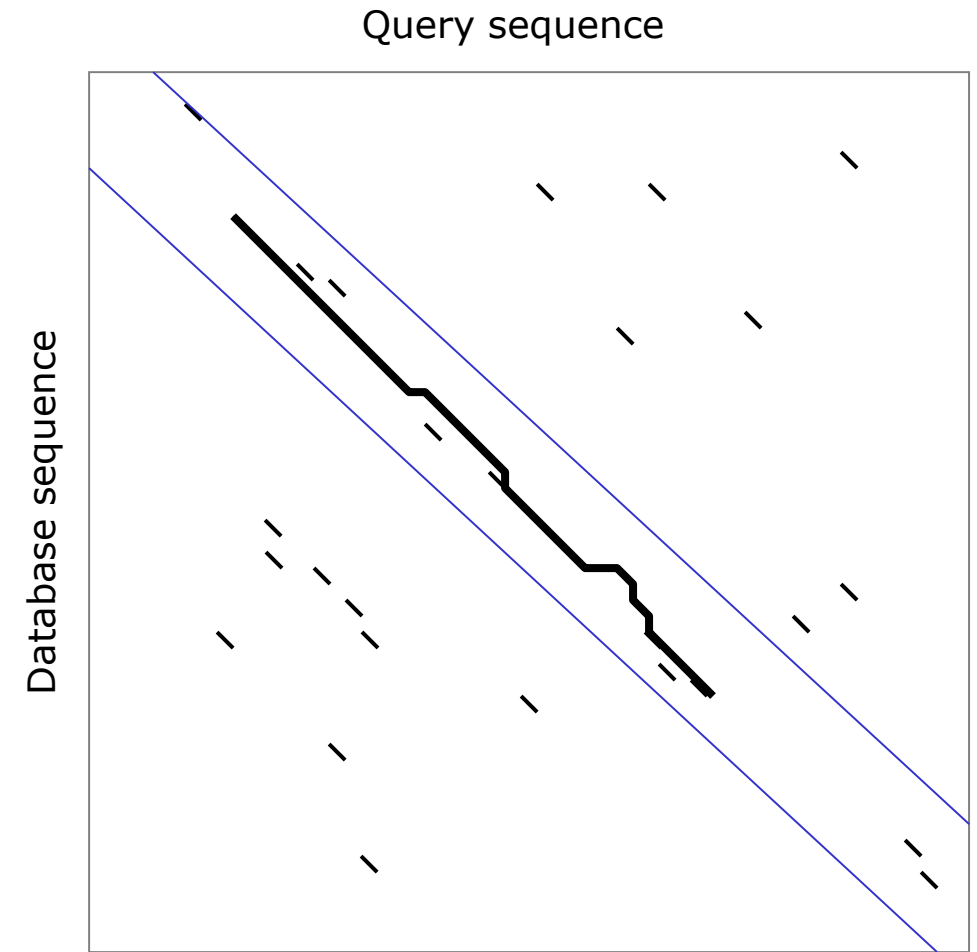


Heurísticas: FASTA

Las palabras de alto score de longitud k son las **semillas** del alineamiento.

- 1) FASTA identifica regiones de **alta densidad** de semillas
- 2) Elige **la region de mayor score**
- 3) Traza una **banda diagonal** que representa los límites para calcular el alineamiento
- 4) Dentro de los límites se calcula el mejor alineamiento usando el algoritmo Smith-Waterman

Pearson WR, Lipman DJ. Improved tools for biological sequence comparison. Proc Natl Acad Sci U S A. 1988 Apr;85(8):2444-8. doi: 10.1073/pnas.85.8.2444. PMID: 3162770; PMCID: PMC280013.



Heurísticas: FASTA

Las palabras de alto score de longitud ktup son las **semillas** del alineamiento.

- 1) FASTA identifica regiones de **alta densidad** de semillas
- 2) Elige **la region de mayor score**
- 3) Traza una **banda diagonal** que representa los límites para calcular el alineamiento
- 4) Dentro de los límites se calcula el mejor alineamiento usando el algoritmo Smith-Waterman

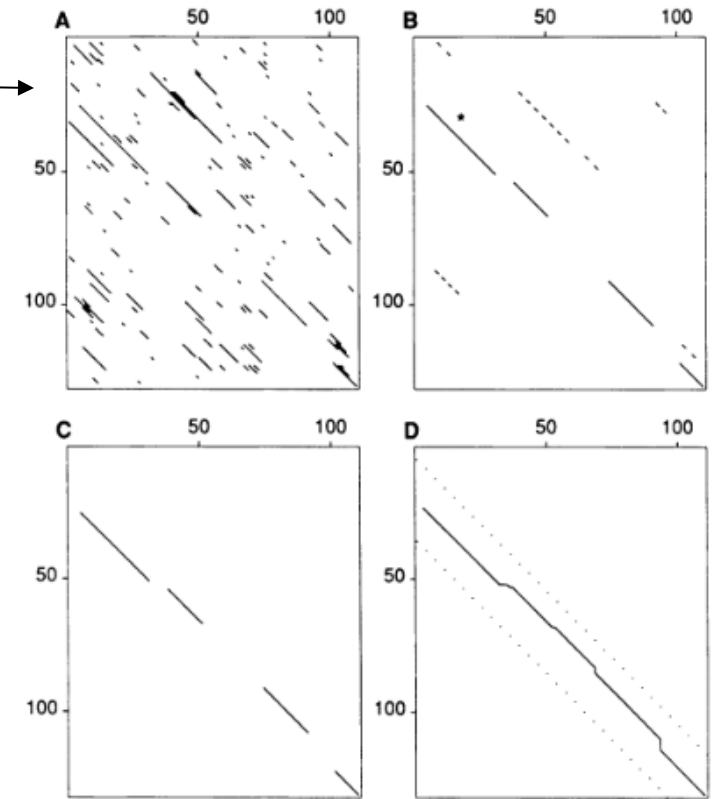


FIG. 1. Identification of sequence similarities by FASTA. The four steps used by the FASTA program to calculate the initial and optimal similarity scores between two sequences are shown. (A) Identify regions of identity. (B) Scan the regions using a scoring matrix and save the best initial regions. Initial regions with scores less than the joining threshold (27) are dashed. The asterisk denotes the highest scoring region reported by FASTP. (C) Optimally join initial regions with scores greater than a threshold. The solid lines denote regions that are joined to make up the optimized initial score. (D) Recalculate an optimized alignment centered around the highest scoring initial region. The dotted lines denote the bounds of the optimized alignment. The result of this alignment is reported as the optimized score.

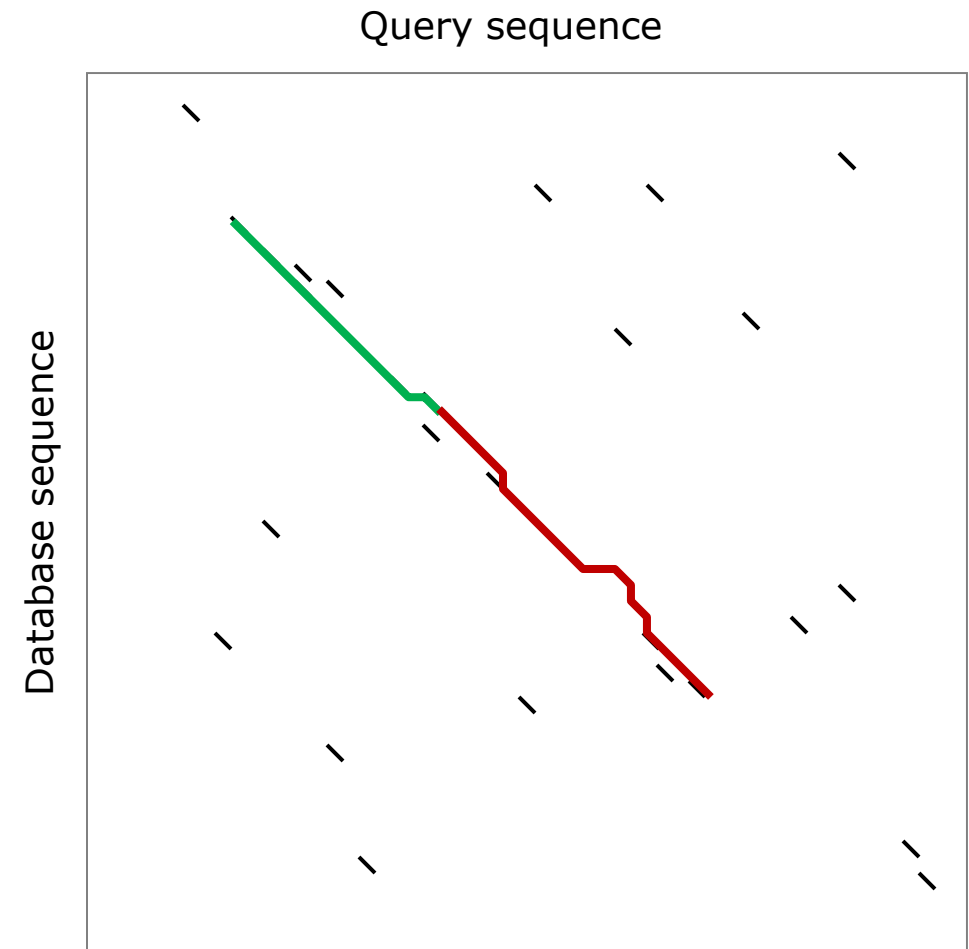
Pearson WR, Lipman DJ. Improved tools for biological sequence comparison. Proc Natl Acad Sci U S A. 1988 Apr;85(8):2444-8. doi: 10.1073/pnas.85.8.2444. PMID: 3162770; PMCID: PMC280013.

Heurísticas: BLAST

Las palabras de alto score de longitud ktup son las **semillas** del alineamiento.

- 1) BLAST identifica regiones de **alta densidad** de semillas
- 2) Regla de proximidad: elige palabras de alto score **a menos de una distancia A**
- 3) Extiende las palabras en ambas direcciones para calcular el alineamiento
- 4) El alineamiento se corta cuando el score cae más que un valor X predeterminado (drop-off)

```
----- (diagonal)
[ Word 1 ] <----- Distance (A) -----> [ Word 2 ]
      |                                     |
      |> (Extension starts here)           |> (And here)
```

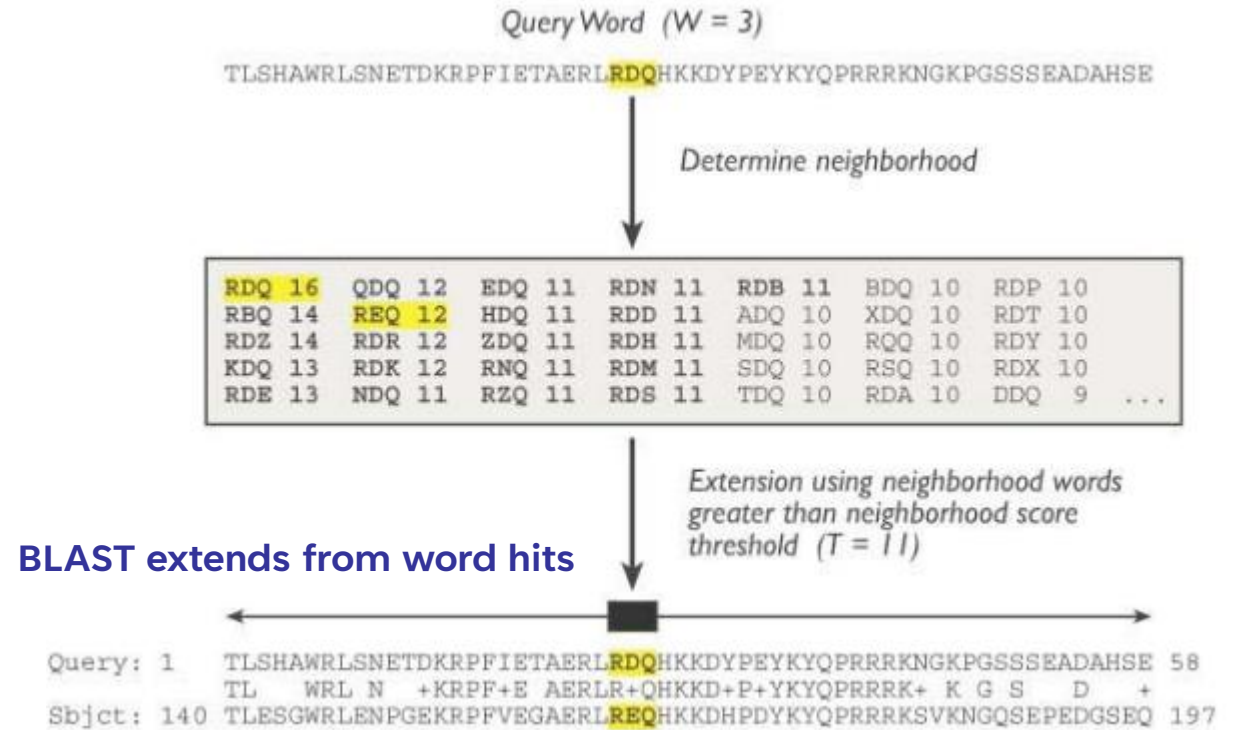


BLAST extends from word hits

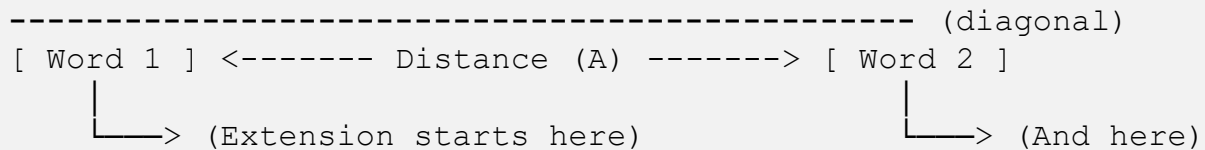
Heurísticas: BLAST

Las palabras de alto score de longitud ktup son las **semillas** del alineamiento.

- 1) BLAST identifica regiones de alta densidad de semillas
- 2) Regla de proximidad: elige palabras de alto score a menos de una distancia A
- 3) BLAST extiende las palabras en ambas direcciones para calcular el alineamiento
- 4) El alineamiento termina cuando el score cae más que un valor X predeterminado (drop-off)

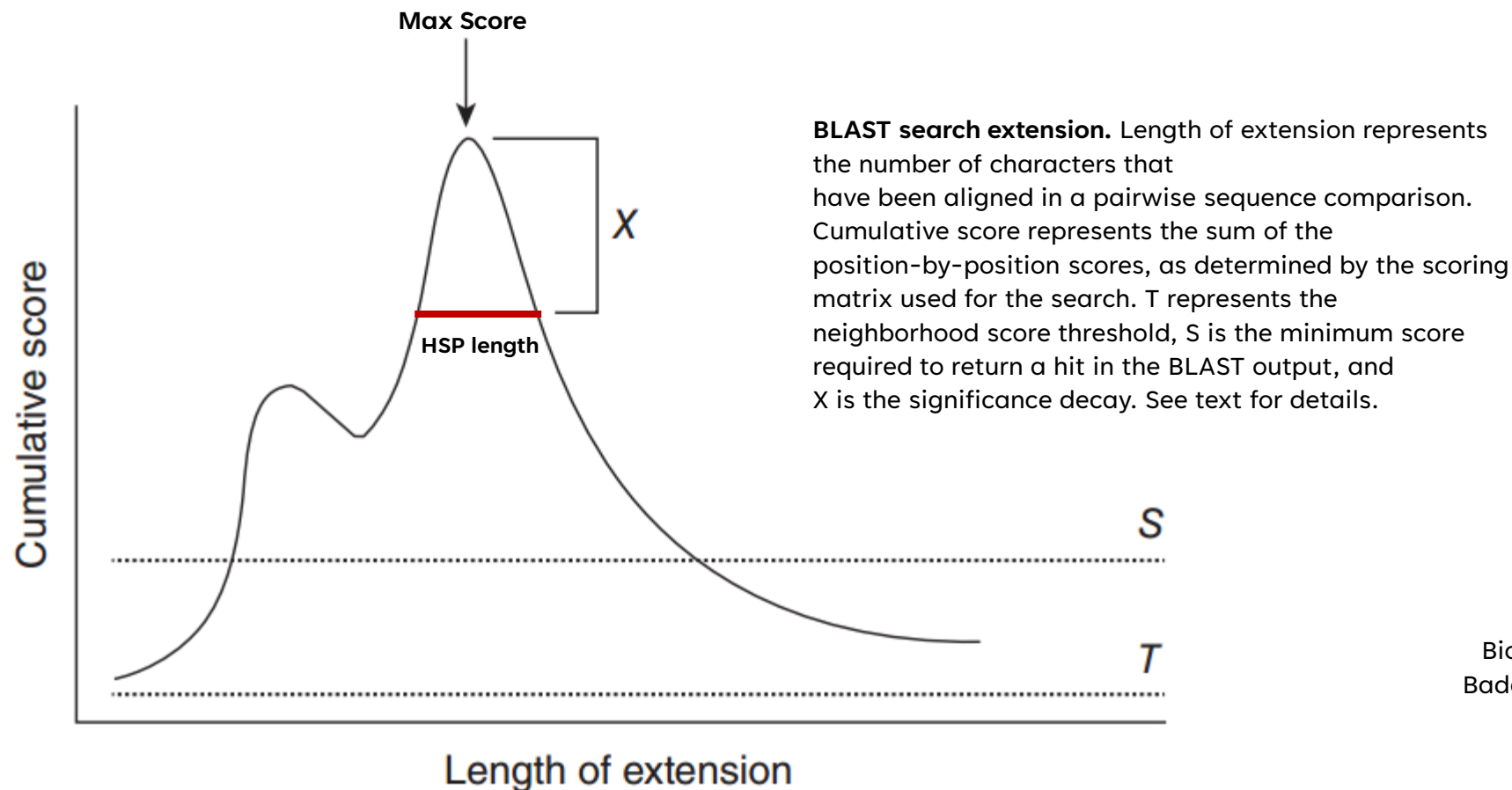


Bioinformatics. Andreas D. Baxevanis, Gary D. Bader, David S. Wishart. John Wiley & Sons, 4th Edition (2020)



BLAST puede producir varios alineamientos (HSPs)

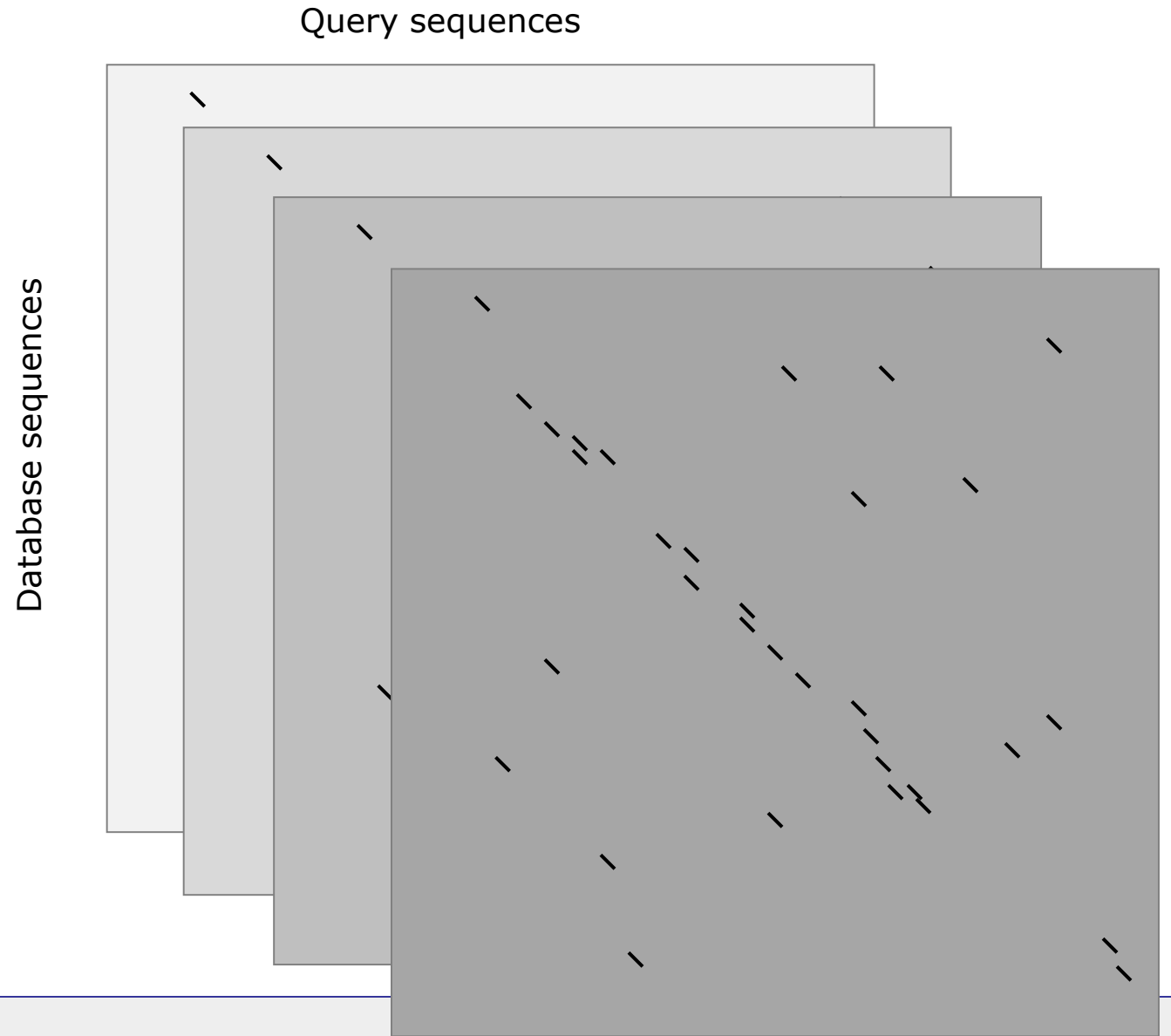
HSP = “*high scoring pair*” (o sea un alineamiento corto, local, de alto score, entre la secuencia “*query*” y una secuencia “*subject*” de la base de datos.)



Bioinformatics. Andreas D. Baxevanis, Gary D. Bader, David S. Wishart. John Wiley & Sons, 4th Edition (2020)

Heurísticas para búsquedas

Esto se repite por cada secuencia de la base de datos (BLAST, FASTA).



I only come to work
for the coffee breaks.



your  cards
someecards.com

INTERVALO

Fin del Slot 1

Búsquedas en bases de datos

Una búsqueda típica tiene 4 elementos básicos:

- 1) Programa / Algoritmo
- 2) Secuencia query
- 3) Base de Datos
- 4) Parámetros opcionales

blastn **blastp** blastx tblastn tblastx

Enter Query Sequence

Enter accession number(s), gi(s), or FASTA sequence(s) [Clear](#)

Query subrange [?](#)

From

To

Or, upload file No file chosen [?](#)

Job Title

Enter a descriptive title for your BLAST search [?](#)

☐ Align two or more sequences [?](#)

Choose Search Set

Database [?](#)

Organism Optional ☐ exclude

Enter organism common name, binomial, or tax id. Only 20 top taxa will be shown. [?](#)

Exclude Optional ☐ Models (XM/XP) ☐ Non-redundant RefSeq proteins (WP) ☐ Uncultured/environmental sample sequences

Program Selection

Algorithm ☒ blastp (protein-protein BLAST)

☐ PSI-BLAST (Position-Specific Iterated BLAST)

☐ PHI-BLAST (Pattern Hit Initiated BLAST)

☐ DELTA-BLAST (Domain Enhanced Lookup Time Acceleration)

Choose a BLAST algorithm [?](#)

Algorithm parameters

General Parameters

Max target sequences [?](#)
Select the maximum number of aligned sequences to display [?](#)

Short queries ☒ Automatically adjust parameters for short input sequences [?](#)

Expect threshold [?](#)

Word size [?](#)

Max matches in a query range [?](#)

Scoring Parameters

Matrix [?](#)

Gap Costs [?](#)

Compositional adjustments [?](#)

Filters and Masking

Filter ☐ Low complexity regions [?](#)

Mask ☐ Mask for lookup table only [?](#)
☐ Mask lower case letters [?](#)

Search database swissprot using Blastp (protein-protein BLAST)
☐ Show results in a new window

```
> blastp myquery swissprot -word_size 5
```

Programa

Secuencia
query

Base de datos

Parámetros
opcionales

https://www.ncbi.nlm.nih.gov/books/NBK279684/table/appendices.T.blastp_application_options/
https://www.ncbi.nlm.nih.gov/books/NBK279684/table/appendices.T.blastn_application_options/

BLAST results

BLAST® » blastp suite » results

◀ Edit Search

Save Search

Descriptions

Graphic Summary

Sequences producing significant alignments

☒ select all 100 sequences selected

- ☒ [RecName: Full=Rhodopsin \[Xenopus laevis\]](#)
- ☒ [RecName: Full=Rhodopsin \[Aequorea victoria\]](#)
- ☒ [RecName: Full=Rhodopsin \[Lithobates pipiens\]](#)
- ☒ [RecName: Full=Rhodopsin \[Rana temporaria\]](#)
- ☒ [RecName: Full=Rhodopsin \[Rhinella marina\]](#)
- ☒ [RecName: Full=Rhodopsin \[Bufo bufo\]](#)
- ☒ [RecName: Full=Rhodopsin \[Ambystoma maculatum\]](#)
- ☒ [RecName: Full=Rhodopsin \[Oryctolagus cuniculus\]](#)
- ☒ [RecName: Full=Rhodopsin \[Gallus gallus\]](#)

Descriptions

Graphic Summary

Alignments

Taxonomy

Alignment view

Pairwise



Restore defaults

Download

100 sequences selected

Download

GenPept Graphics

Next Previous Descriptions

RecName: Full=Rhodopsin [Xenopus laevis]

Sequence ID: [P29403.1](#) Length: 354 Number of Matches: 1

Range 1: 1 to 354 GenPept Graphics

Next Match Previous Match

Score	Expect	Method	Identities	Positives	Gaps
724 bits(1870)	0.0	Compositional matrix adjust.	351/354(99%)	352/354(99%)	0/354(0%)
Query 1		MNGTEGPNFYVPM SNK TGVVRS PFDYPQYYLAEPWQYSALAAYMFLLILLGLPINFMTLF			60
Sbjct 1		MNGTEGPNFYVPM SNK TGVVRS PFDYPQYYLAEPWQYSALAAYMFLLILLGLPINFMTLF			60
Query 61		VTIQHKLRTP LNYILLNLVFANHF MVLCGFTVTMYTSMHGYFIFGQTGCYIEGFFATLG			120
Sbjct 61		VTIQHKLRTP LNYILLNLVFANHF MVLCGFTVTMYTSMHGYFIFGQTGCYIEGFFATLG			120
Query 121		GEVALWSLVVLAVERYMVVCKPMANFRFGENHAIMGVAF TWIMALSCAAPPLFGWSRYIP			180
Sbjct 121		GEVALWSLVVLAVERYMVVCKPMANFRFGENHAIMGVAF TWIMALSCAAPPLFGWSRYIP			180
Query 181		EGMQSCGVDYYTLKPEVNNE SFVIYMFIVHFTIPLIVIFFCYGRLLCTVKEAAQQQES			240
Sbjct 181		EGMQSCGVDYYTLKPEVNNE SFVIYMFIVHFTIPLIVIFFCYGRLLCTVKEAAQQQES			240
Query 241		ATTQKAEKEVTRMVVIMVVFLLICWVPYAYVAFYIFTHQGSNFGPVFMTVPAFFAKSSAI			300
Sbjct 241		ATTQKAEKEVTRMVVIMVVFLLICWVPYAYVAFYIFTHQGSNFGPVFMTVPAFFAKSSAI			300
Query 301		YNPVIYIVLNKQFRNCLITTLCCGKNPFGDEDGSSAATSKTEASSVSSSQVSPA			354
Sbjct 301		YNPVIYIVLNKQFRNCLITTLCCGKNPFGDEDGSSAATSKTEASSVSSSQVSPA			354

Related Information

[Gene](#) - associated gene details
[AlphaFold Structure](#) - 3D structure displays

BLASTP results

> **blastp** **myquery** **swissprot** **-word_size** **5**

Programa Secuencia query Base de datos Parámetros opcionales

```
>O62791.1 RecName: Full=Rhodopsin [Delphinus delphis]
Length=348

Score = 616 bits (1589), Expect = 0.0, Method: Compositional matrix adjust.
Identities = 287/354 (81%), Positives = 326/354 (92%), Gaps = 6/354 (2%)

Query 1 MNGTEGPNFYVPM SNKTGVVRSPFDYPQYYLAEPWQYSALAAYMFLLILLGLPINFMTLF 60
Sbjct 1 MNGTEG NFYVP SNKTGVVRSPF+YPQYYLAEPWQ+S LAAYMFLLI+LG PINF+TL+
MNGTEGLNFYVFP SNKTGVVRSPFEYPQYYLAEPWQFSVLAAYMFLLIVLGFPINFLTLY 60

Query 61 VTIQHKKLRTPLNYILLNLVFNHFVLCGFTVTMYTSMHGYFIFGQTGCYIEGFFATLG 120
Sbjct 61 VT+QHKKLRTPLNYILLNL AN FMV GFT T+YTS+H YF+FG TGC +EGFFATLG
VTVQHKKLRTPLNYILLNLAVANLFMVFGGFTTTLYTSLHAYFVFGPTGCNLEGFFATLG 120

Query 121 GEVALWSLVVLAVERYMVVCKPMANFRFGENHAIMGVAFTWIMALSCAAPPLFGWSRYIP 180
Sbjct 121 GEIALWSLVVLAIERVYVVKPMSNFRFGENHAIMGLALTWIMAMACAAPPLVGWSRYIP 180

Query 181 EGMQCSCGVDYYTLKPEVNNESEFVIYMFIVHFTIPLIVIFFCYGRLLCTVKEAAAQQQES 240
Sbjct 181 EGMQCSCG+DYITL PEVNNESEFVIYMF+VHFTIPL++IFFCYG+L+ TVKEAAAQQQES
EGMQCSCGIDYYTLSPEVNNESEFVIYMFVVHFTIPLVLIIFFCYGQLVFTVKEAAAQQQES 240

Query 241 ATTQKAEKEVTRMVVIMVVFFLICWVPYAYVAFYIFTHQGSNFGPVFMTVPAFFAKSSAI 300
Sbjct 241 ATTQKAEKEVTRMV+IMVV FLICWVPYA VAFYIFTHQGS+FGP+FMT+P+FFAKSS+I
ATTQKAEKEVTRMVIIMVVAFLICWVPYASVAFYIFTHQGSDFGPIFMTIPSFFAKSSSI 300

Query 301 YNPVIYIVLNKQFRNCLITTLCCGKNPFGDEGSSAATSKTEASSVSSSQVSPA 354
Sbjct 301 YNPVIYI++NKQFRNC++TTLCCG+NP GD++ S+ A SKTE +SQV+PA
YNPVIYIMMNKQFRNCMLTTLCCGRNPLGDDEASTTA-SKTE-----TSQVAPA 348
```

BLASTP 2.12.0+

Reference: Stephen F. Altschul, Thomas L. Madden, Alejandro A. Schaffer, Jinghui Zhang, Zheng Zhang, Webb Miller, and David J. Lipman (1997), "Gapped BLAST and PSI-BLAST: a new generation of protein database search programs", Nucleic Acids Res. 25:3389-3402.

Reference for composition-based statistics: Alejandro A. Schaffer, L. Aravind, Thomas L. Madden, Sergei Shavirin, John L. Spouge, Yuri I. Wolf, Eugene V. Koonin, and Stephen F. Altschul (2001), "Improving the accuracy of PSI-BLAST protein database searches with composition-based statistics and other refinements", Nucleic Acids Res. 29:2994-3005.

Database: Non-redundant UniProtKB/SwissProt sequences
487,492 sequences; 185,956,103 total letters

Query= L07770

Length=354

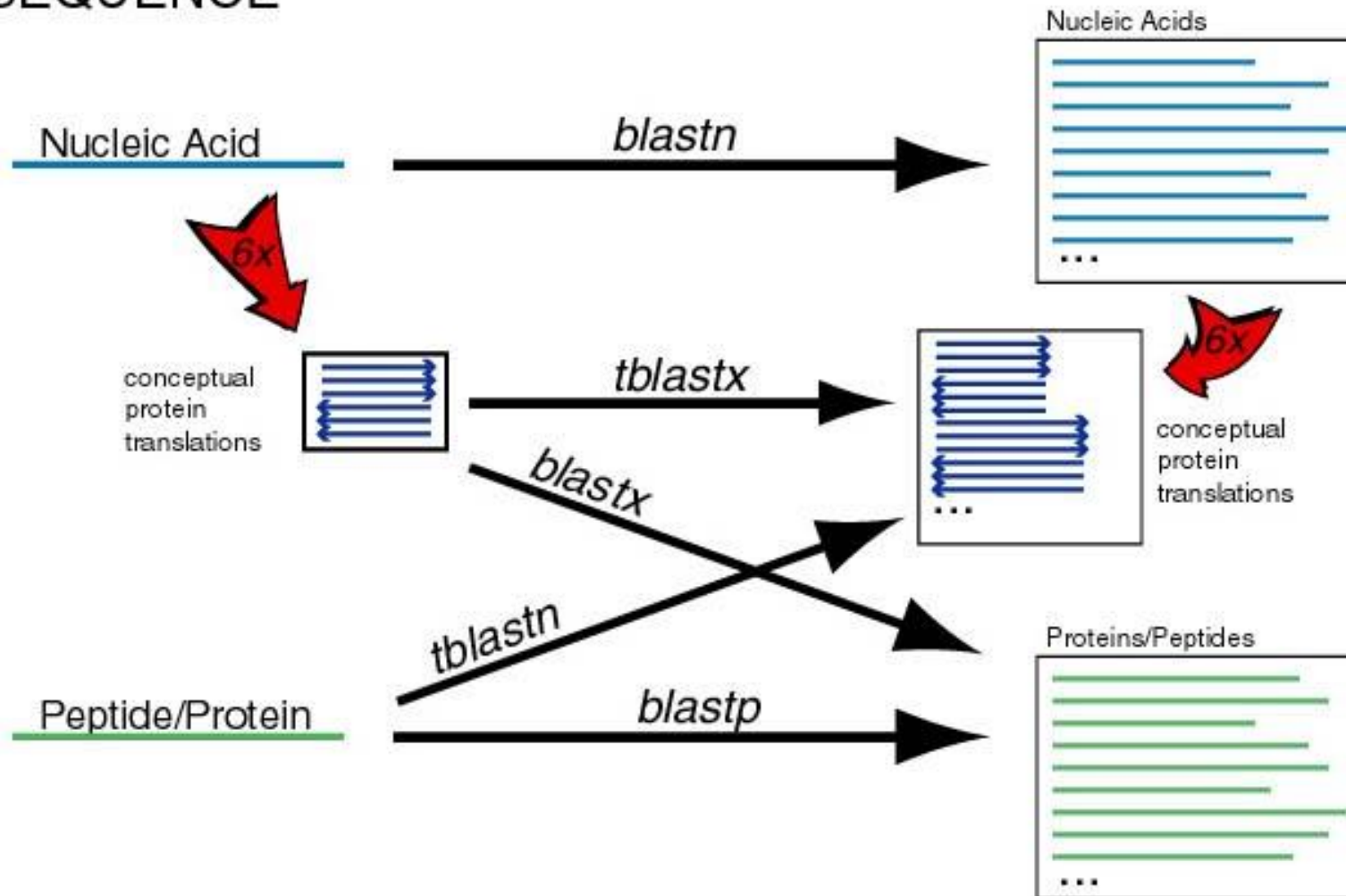
Sequences producing significant alignments:

	Score (Bits)	E Value
P29403.1 RecName: Full=Rhodopsin [Xenopus laevis]	724	0.0
P51470.1 RecName: Full=Rhodopsin [Aequipecten irradians]	694	0.0
P31355.1 RecName: Full=Rhodopsin [Lithobates pipiens]	688	0.0
P56516.1 RecName: Full=Rhodopsin [Rana temporaria]	681	0.0
P56515.1 RecName: Full=Rhodopsin [Rhinella marina]	672	0.0
P56514.1 RecName: Full=Rhodopsin [Bufo bufo]	660	0.0
Q90245.1 RecName: Full=Rhodopsin [Ambystoma tigrinum]	654	0.0
P49912.1 RecName: Full=Rhodopsin [Oryctolagus cuniculus]	630	0.0
P22328.1 RecName: Full=Rhodopsin [Gallus gallus]	628	0.0
Q68J47.1 RecName: Full=Rhodopsin [Loxodonta africana]	628	0.0
Q95KU1.1 RecName: Full=Rhodopsin [Felis catus]	626	0.0
P15409.2 RecName: Full=Rhodopsin [Mus musculus]	625	0.0
P52202.1 RecName: Full=Rhodopsin [Alligator mississippiensis]	622	0.0
Q28886.1 RecName: Full=Rhodopsin [Macaca fascicularis]	621	0.0

BLAST: algoritmos

QUERY SEQUENCE

DATABASE



- Qué hacemos cuando estamos comparando dos secuencias que no son similares, pero que muestran un alineamiento prometedor?
- Necesitamos un **test de significancia**
- Tenemos que responder a la pregunta:
 - Cuál es la probabilidad de que un alineamiento similar (con un score similar) ocurra entre proteínas no relacionadas?

Solamente con el score basta?

Dos secuencias **no relacionadas** (que no son homólogas), pueden tener alineamientos con **el mismo puntaje** o muy cercano?

Dos estrategias alternativas:

- Generar **secuencias al azar** de la misma longitud y composición que la secuencia query y alinearlas
 - Karlin & Altschul (1990); Altschul et al (1994); Altschul & Gish (1996)
 - Suponer que la base de datos **es suficientemente grande** y está compuesta mayormente por secuencias **no relacionadas**
 - Pearson (1998)
-
- Analizar **la distribución de scores** que se obtiene

<https://www.ncbi.nlm.nih.gov/BLAST/tutorial/Altschul-1.html>

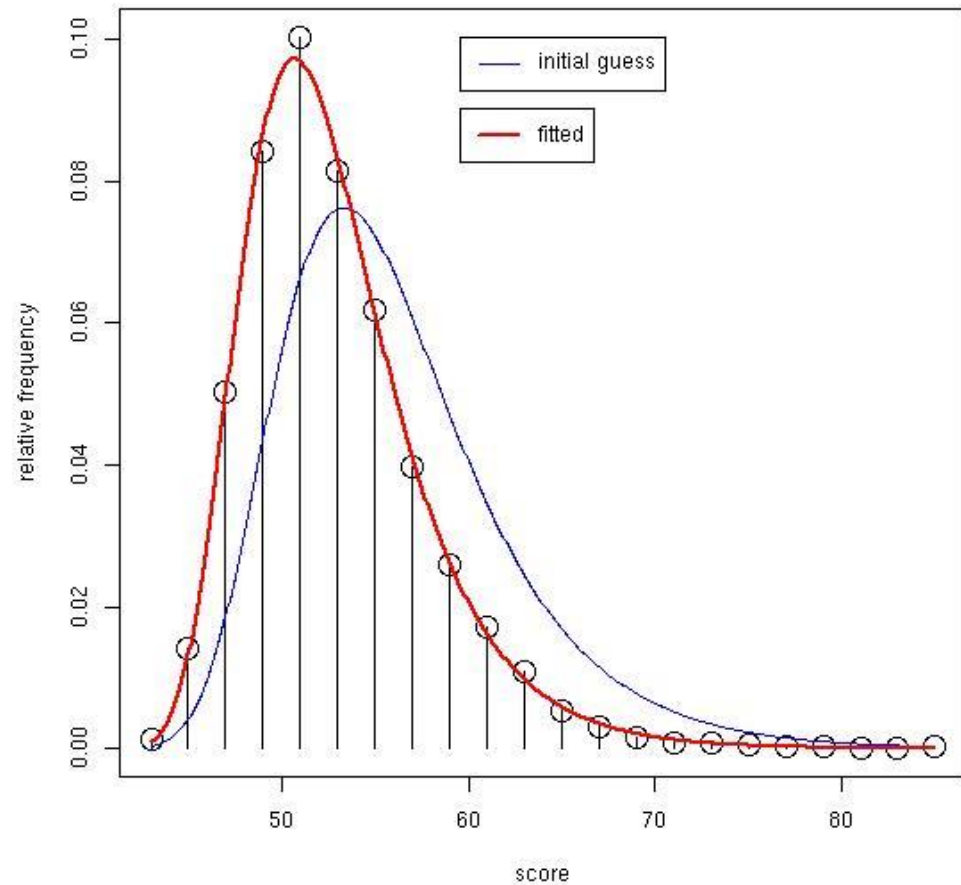
Pearson WR. Empirical statistical estimates for sequence similarity searches. J Mol Biol. 1998 Feb 13;276(1):71-84. doi: 10.1006/jmbi.1997.1525. PMID: 9514730.

Karlin S, Altschul SF. Methods for assessing the statistical significance of molecular sequence features by using general scoring schemes. Proc Natl Acad Sci U S A. 1990 Mar;87(6):2264-8. doi: 10.1073/pnas.87.6.2264. PMID: 2315319; PMCID: PMC53667.

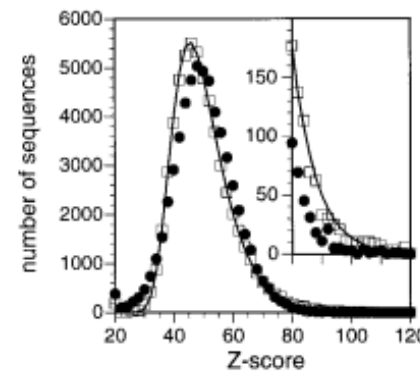
Análisis de las distribuciones de scores

In a database search (BLAST/FASTA) the alignment scores **do not** follow a normal/Gaussian distribution!

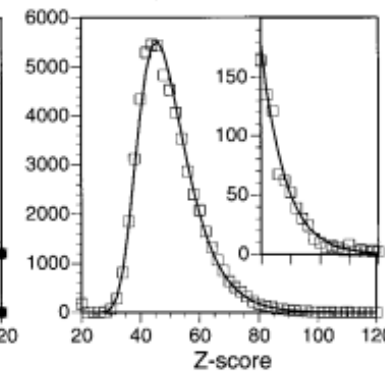
Extreme Value Distribution, Empirical method



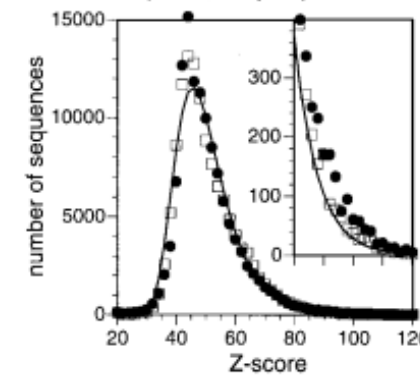
A. Smith-Waterman



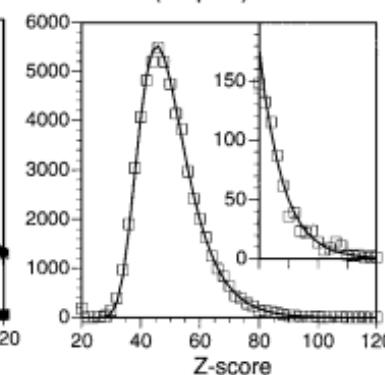
B. FASTA (ktup=2)



C. FASTA (DNA, ktup=4)



D. FASTX (ktup=2)



The Gumbel/Extreme value distribution

In a database search (BLAST/FASTA) the alignment scores follow a Gumbel or Extreme Value distribution!

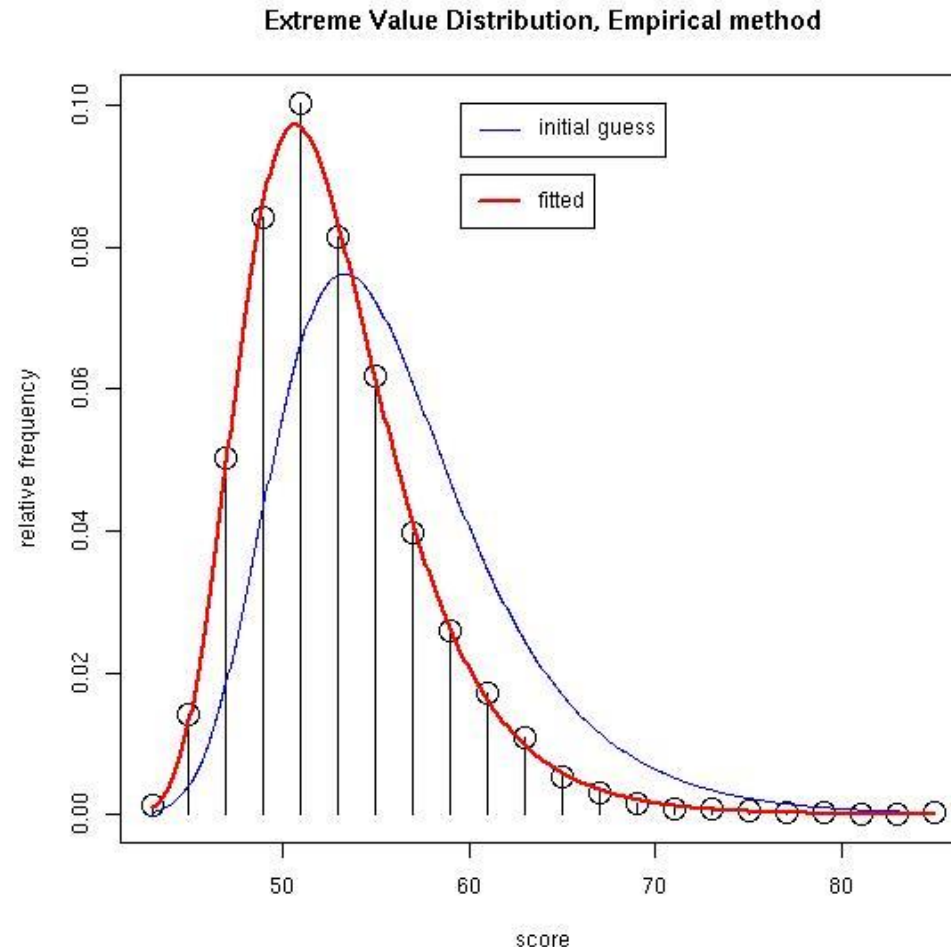
$$E = K m n e^{-\lambda S}$$

E is the **number** of alignments with a score = S

m, n : length of the sequences

K, λ : estimated parameters estimated
(depend on the scoring matrix and the size of the database)

e : Euler's number (mathematical constant)



Los hits pueden ser ordenados de acuerdo a su E-value o a su Score.

El E-value – más conocido como **EXPECT** value – es una función del score, el tamaño de la base de datos y de la longitud de la secuencia 'query'.

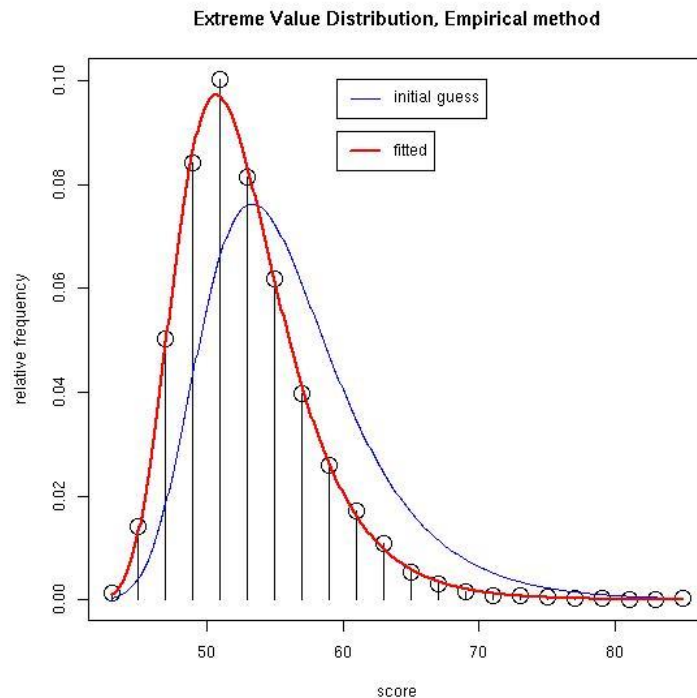
E-value: Número de alineamientos con un score = S que se espera encontrar si la base de datos es una colección de letras al azar.

Ejemplo: En el caso de un score=1 (un match o identidad) debería haber un número enorme de alineamientos. Uno espera encontrar menos alineamientos con un score de 5, 10, etc. Eventualmente, cuando el score es lo suficientemente alto, uno espera encontrar un número insignificante de alineamientos que sean debidos al azar.

Valores de E-value menores que $1e-6$ ($1 * 10^{-6}$) son generalmente muy buenos para proteínas, mientras que $E < 1e-2$ puede considerarse significativo. Es posible que un hit cuyo $E > 1$ sea biológicamente importante, aunque es necesario analizarlo más detalladamente para confirmarlo.

Observed vs expected

Si la base de datos es suficientemente grande y contiene mayoritariamente secuencias no relacionadas la distribución de scores **observados** debería coincidir bastante con la distribución de scores **esperados por azar** (Pearson 1998)



Pearson WR. Empirical statistical estimates for sequence similarity searches. J Mol Biol. 1998 Feb 13;276(1):71-84. doi: 10.1006/jmbi.1997.1525. PMID: 9514730.

Karlin S, Altschul SF. Methods for assessing the statistical significance of molecular sequence features by using general scoring schemes. Proc Natl Acad Sci U S A. 1990 Mar;87(6):2264-8. doi: 10.1073/pnas.87.6.2264. PMID: 2315319; PMCID: PMC53667.

$$E(S > x) = p(S > x) D$$

El número de alineamientos con un score $> S$ se incrementa linealmente con el tamaño de la base de datos

Tamaño de la base de datos: un ejemplo

Objetivo: encontrar en *E. coli* el gen homólogo al gen AroA de *Bacillus subtilis* (DAHP synthase)

- **Database = *E. coli* proteome (4,283 proteínas)**
 - **kdsA, $E(4283) < 0.00015$**
- **Database = Swissprot (74,417 proteínas)**
 - **kdsA, $E(74417) < 0.0017$**
- **Database = OWL (260,784 proteínas)**
 - **kdsA, $E(260784) < 0.0085$**
- **El mismo alineamiento, con el mismo score es 50 veces más significativo en la base de datos más chica.**

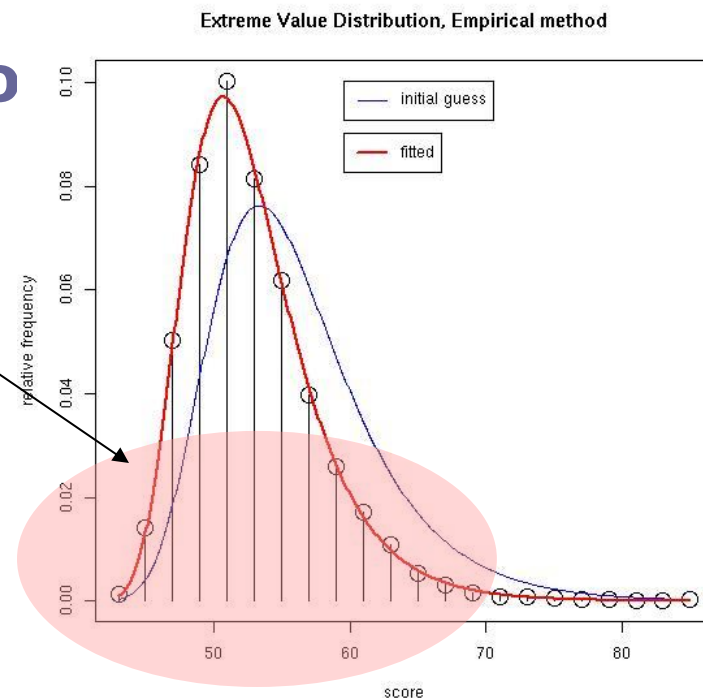
Disclaimer: los tamaños de las bases de datos son ilustrativos y corresponden a 1998-1999 cuando se hizo este ejercicio.

Bleasby AJ, Akrigg D, Attwood TK. OWL--a non-redundant composite protein sequence database. Nucleic Acids Res. 1994 Sep;22(17):3574-7. PMID: 7937061; PMCID: PMC308323.

Problema: estamos buscando homólogos lejanos

Estrategia:

- Buscar en bases de datos **pequeñas primero**
- Repetir la búsqueda en una base de datos pequeña con un **algoritmo más sensible** (`word_size / ktup = 1` o un algoritmo exacto sin heurísticas, ej `ssearch`)
- Dudar de alineamientos de bajo score



Límites de la estadística

- En ciertos casos, la estadística de los alineamientos falla
 - Lo que falla son las suposiciones que hicimos para llegar al modelo estadístico
- Se obtienen estimaciones incorrectas del E-value cuando
 - Se usan penalidades de gap incorrectas o extremas
 - Existen regiones de **baja complejidad** en la secuencia query
 - La secuencia query tiene una **composición sesgada**

>sp|P42858|HD_HUMAN Huntingtin OS=Homo sapiens OX=9606 GN=HTT
MATLEKLMKAFESLKSF**QQQQQQQQQQQQQQQQQQQQQQPPPPPPPPPP**QLPQPPPPQAQP
LLPQPQ**PPPPPPPPPP**GPAVAEEPLHRPKKELSATKKDRVNHCLTICENIVAQSVRNSPE
FQKLLGIAMELFLLCSDDAESDVRMVADECLNKVIKALMDSNLPRLQLELYKEIKKNGAP
RSLRAALWRFAELAHLVRPQKCRPYLVNLLPCLTRTSKRPEESVQETLAAAVPKIMASFG
NFANDNEIKVLLKAFIANLKSSSPTIRRTAAGSAVSICQHSRRTQYFYSWLLNVLLGLLV
PVEDEHSTLLILGVLLTLRYLVPLLQQQVKDTSLKGSFGVTRKEMEVSPPSAEQLVQVYEL
TLHHTQHQDHNVTGALELLQQLFRTPPPELLQTLTAVGGIGQLTAAKEESGGRSRSGSI
[...]

>sp|P04156|PRIO_HUMAN Major prion protein OS=Homo sapiens OX=9606 GN=PRNP
MANL**GC**WMLVLFVATWSDL**GL**CKKRPKP**GG**WNT**GG**SRYP**GQ**SP**GG**NRYPPQ**GGGG**WGQP
H**GGGG**WGQPH**GGGG**WGQPH**GGGG**WGQPH**GGGG**WGQ**GGG**THSQWNKPSKPKTNMKHMAGAAAA**GA**
VV**GG**L**GG**YML**GS**AMSRPIIH**FG**SDYEDRYYYRENMHRYPNQVYYRPMDEYSNQNNFVHDCV
NITIKQHTVTTTTTK**GEN**FTETDVKMMERVVEQMCITQYERESQAYYQR**G**SSMVLFSPPV
ILLISFLIFLIV**G**

Evaluando la estadística en casos extremos

```

opt      E()
< 20    13      0:=
22      0      0:      one = represents 22 library sequences
24      0      0:
26      0      0:
28      1      3:~
30     11     19:~
32     46     75:====~
34    242    204:=====+~
36     493    419:=====+~
38     788    692:=====+~
40    1055    965:=====+~
42    1275   1180:=====+~
44    1299   1302:=====+~
46    1251   1326:=====+~
48    1186   1269:=====+~
50    1077   1158:=====+~
52     907   1018:=====+~
54     849    870:=====+~
56     714    727:=====+~
58     570    596:=====+~
60     456    483:=====+~
62     393    387:=====+~
64     313    308:=====+~
66     268    243:=====+~
68     219    192:=====+~
70     191    150:=====+~
72     127    117:=====+~
74      93     91:=====+~
76      91     71:=====+~
78      44     55:=====+~
80      33     43:=====+~
82      22     33:=====+~
84      32     26:=====+~
86      19     20:=====+~
88      19     16:~
90       8     12:~
92       8      9:~
94       5      7:~
96       2      6:~
98       3      4:~
100      1      3:~
102      3      3:~
104      0      2:~
106      1      2:~
108      0      1:~
110      0      1:~
112      0      1:~
114      0      1:~
116      0      0:~
118      1      0:=
>120     7      0:=

inset = represents 1 library sequences

```

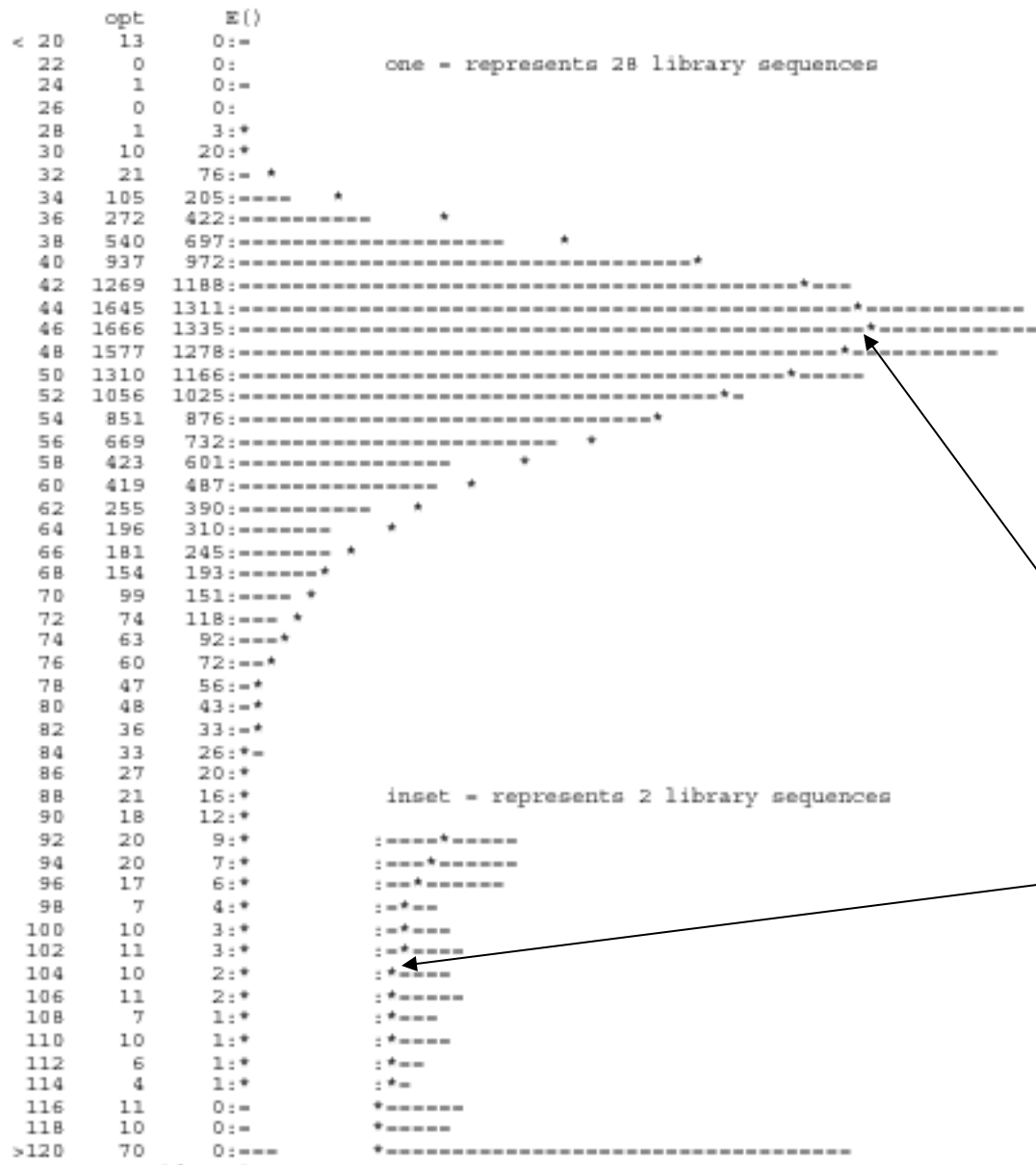
Distribución (histograma) de scores esperados y observados

* $\rightarrow$ observados

= $\rightarrow$ esperados

Mirar el E-value de la
secuencia no relacionada
con mayor score

Evaluando la estadística (cont)



Si los histogramas de Observados vs Esperados coinciden

Y si el E del mejor alineamiento no relacionado es <<< 1

La estimaciones estadísticas están funcionando bien.

Y en este caso?

* → observados
= → esperados

Buscando homólogos en los límites

- **Secuencias homólogas distantes a menudo no tienen similitud estadísticamente significativa**
- **Secuencias con regiones de baja complejidad pueden tener similitud estadísticamente significativas, aunque no sean homólogas**
- **Secuencias homólogas generalmente son similares sobre toda la longitud de la secuencia o de un dominio**
- **Secuencias homólogas comparten un ancestro común**
 - **Si hay homología entre A y B; entre B y C; y entre C y D, A y D deben ser homólogos, aun cuando no muestren similitud estadísticamente significativa**

Low complexity sequences

- **Secuencias (o sub-secuencias) con bajo contenido de información**
 - AAAAAAAAAAAAAAAAAAAAAA
 - CAACAACAACAACAACAA
- **Secuencias con sesgo en la composición de bases (nucleótidos) o residuos (aminoácidos)**
- **Por el bajo contenido de información y el sesgo en la composición, suelen dar falsos positivos en las búsquedas por similitud**
 - PolyQ: proteínas con trectos de polyglutamina no están relacionadas por ancestría
 - Ej OTX2 (Transcription factor); CREB-binding protein (connects proteins with different functions); MED15/GAL11 (Subunit of the RNA polymerase II mediator complex)

BLAST compositional adjustment of scoring matrices

Las matrices de scoring (ej BLOSUM62) son derivadas a partir de alineamientos de proteínas globulares, con una **composición de aminoácidos definida**

Pero muchas veces las secuencias “query” tienen composiciones de aminoácidos con sesgos muy marcados y diferentes a las de las secuencias utilizadas para derivar las matrices

Ej proteínas hidrofóbicas, Cys-rich, proteínas codificadas por genomas con alto sesgo (AT rich, GC rich), que afectan los codones más utilizados

En esos casos hay maneras de ajustar las matrices (ej recalcular BLOSUM62 on-the-fly) para que funcionen mejor frente a estos casos

BLAST compositional adjustment of scoring matrices

Query = human insulin NP_000198
Program = blastp
Database = *C. elegans* RefSeq

Option = NO compositional adjustment

```
>[ref|NP_501926.1| UG INSulin related family member (ins-1) [Caenorhabditis elegans]
Length=109

Score = 34.7 bits (78), Expect = 0.009
Identities = 30/100 (30%), Positives = 41/100 (41%), Gaps = 14/100 (14%)

Query 11  LALLALWGPDPAAAFVNQHLGSHLVEALYLVCGERGFFYTPKTRREAEDLQVGQVELGG 70
          LA+L L P P+ A + LCGS L L VC + +R A+
Sbjct 17  LAILLSSPTPSDASIR--LCGSRLLLLAVCRNQLCTGLTAFKRSADQSY----- 66

Query 71  GPGAGSLQPLALEGSLQKRG-IVEQCCTSLCSLYQLENYC 109
          A + + L QKRG I +CC CS L+ +C
Sbjct 67  ---APTTRDLFHIHQKRGGIATECCEKCSFAYLKTFC 103
```

Scoring Parameters

Matrix

BLOSUM62

Gap Costs

Existence: 11 Extension: 1

Compositional adjustments

Conditional compositional score matrix adjustment



Option = conditional compositional score matrix adjustment

```
>[ref|NP_501926.1| UG INSulin related family member (ins-1) [Caenorhabditis elegans]
Length=109

Score = 33.5 bits (75), Expect = 0.020, Method: Compositional matrix adjust.
Identities = 27/100 (27%), Positives = 39/100 (39%), Gaps = 12/100 (12%)

Query 10  LLALLALWGPDPAAAFVNQHLGSHLVEALYLVCGERGFFYTPKTRREAEDLQVGQVELG 69
          LA+L L P P+ A + LCGS L L VC + +R A+
Sbjct 16  FLAILLSSPTPSDASIR--LCGSRLLLLAVCRNQLCTGLTAFKRSADQ-----S 65

Query 70  GPGAGSLQPLALEGSLQKRGIVEQCCTSLCSLYQLENYC 109
          P L + ++ GI +CC CS L+ +C
Sbjct 66  YAPTTRDL--FHIHQKRGGIATECCEKCSFAYLKTFC 103
```

BLAST compositional adjustment of scoring matrices

Fig. 1 Examples of alignment extensions yielded by compositional adjustment of the scoring system.

(a)

```
78 ISGAILFEETL FQKNEAGVPMVNLLHNENIIPGIKVDKGLVNIPCTDEE--KSTQGLDGLAERCKEYYKAGA 147
I GAILFE+T+ K ++ L + ++P +K+DKGL ++ + K L L +R E + G
67 ILGAILFEQTMD SKIDGKYTADFLWEEKKVL PFLKIDKGLNDLDADGVQTMKPNPTLADLLKRANERHIFG- 137

148 RFAKWRTVLVIDTAKGKPTDLS-IHETAWGLARYASICQQNRLVPIVEPEI 197
K R+V+ K+ P ++ + E + +A A + L+PI+EPE+
138 --TKMRSVI----KKASPAGIARVVEQQFEVA--AQVVAAG-LIPIIEPEV 179
```

(b)

```
78 ISGAILFEETL FQKNEAGVPMVNLLHNENIIPGIKVDKGLVNIPCTDEE--KSTQGLDGLAERCKEYYKAGA 147
I GAILFE+T+ K + L + + ++P +K+DKGL ++ + + K L L +R+ E + G
67 ILGAILFEQTMD SKIDGKYTADFLWEEKKVL PFLKIDKGLNDLDADGVQTMKPNPTLADLLKRANERHIFG- 137

148 RFAKWRTVLVIDTAKGKPTDLS-IHETAWGLARYASICQQNRLVPIVEPEILADGPHSIEVCAVVTQKVLSC 218
K+R+V+ K+ P ++ + E + +A A ++ L+PI+EPE+ ++ ++ C + + +
138 --TKMRSVI----KKASPAGIARVVEQQFEVA--AQVVAAG-LIPIIEPEVDINNVDKVQ-CEEILRDEIRK 199

219 VFKALQE-NGVLLEGALLKPNMVTAGYECTAKTTTQDVGFLTVRTLRRTVPPALPGVVFLSGGQSEEEAS 287
+ AL E ++V+L+ L P + E T P + VV LSGG S E+A+
200 HLNALPETS NVMLKLTL--PTVENLYEEFTKH-----PRVVRVVALSGGYSREKAN 248

219 VFKALQE-NGVLLEGALLKPNMVTAGYECTAKTTTQDVGFLTVRTLRRTVPPALPGVVFLSGGQSEEEAS 287
+ + ++V+L+ L P + + E T P + VV LSGG S E+A+
201 LNALPETS NVMLKLTL--PTVENLYEEFTKH-----PRVVRVVALSGGYSREKAN 248
```

(c)

```
78 ISGAILFEETL FQKNEAGVPMVNLLHNENIIPGIKVDKGLVNIPCTDEE--KSTQGLDGLAERCKEYYKAGA 147
I GAILFE+T+ K + L + + ++P +K+DKGL ++ + + K + L L +R+ E + G
67 ILGAILFEQTMD SKIDGKYTADFLWEEKKVL PFLKIDKGLNDLDADGVQTMKPNPTLADLLKRANERHIFG- 137

148 RFAKWRTVLVIDTAKGKPTDLS-IHETAWGLARYASICQQNRLVPIVEPEILADGPHSIEVCAVVTQKVLSC 218
K+R+V+ K+ P ++ + E + +A A ++ L+PI+EPE+ ++ ++ ++ +++
138 --TKMRSVI----KKASPAGIARVVEQQFEVA--AQVVAAG-LIPIIEPEVDINNVDKVQCEEILRDEIRKH 200

219 VFKALQENGVLLEGALLKPNMVTAGYECTAKTTTQDVGFLTVRTLRRTVPPALPGVVFLSGGQSEEEAS 287
+ + ++V+L+ L P + + E T P + VV LSGG S E+A+
201 LNALPETS NVMLKLTL--PTVENLYEEFTKH-----PRVVRVVALSGGYSREKAN 248
```

a) standard BLOSUM-62 substitution matrix

b) composition-adjusted matrix derived from BLOSUM-62 with unconstrained relative entropy

c) composition-adjusted matrix derived from BLOSUM-62 with relative entropy constrained to 0.566 bits

