

Sequence Alignment & Similarity Search.

Motivations:

1. A new gene $\rightarrow$ what is its function?

Compare with the database of known genes
& based on similarity, formulate hypothesis.

e.g. 1984 $\rightarrow$ newly discovered cancer-causing
v-src oncogene (resides in a virus) with all known genes & found
a close match with a gene involved in
growth & development of cell.

2. A segment of chromosome identified for
certain function $\rightarrow$ what
gene resides there? $\Rightarrow$ Compare sequenced
segment with the gene database

e.g. cystic fibrosis gene in chromosome 7 (human)
in 1989, matches with a gene
responsible for coding ATP binding proteins.

⑪ DNA "sequence similarity" is more than a "string matching" problem.

- Evolutionary processes: Mutation, replication error causes insertion & deletion & substitution.
- Exon-intron chains.

e.g.

ATATAT
TATATA

→ Similar from the angle of genome analysis.

⊛ Edit distance: (Vladimir Levenshtein, 1966)

"Edit distance" between two strings is the "minimum no." of "editing operations" needed to transform from one string into another, where the edit operations are insertion of a symbol, deletion of a symbol & substitution of one symbol for another.

e.g.

TGCATAT
↓ delete T
TGCATATA
↓ delete A
TGCATA
↓ insert A at the front
ATGCATA
↓ substitute C for G in the 3rd. pos.
ATCCAT
↓ insert a G before the last A
ATCCGAT

∴ edit dist. (TGCATAT, ATCCGAT) = 5

TGCATAT

↓ insert A at the front
ATGCATAT
↓ delete T in the 6th pos.
ATGC AAT
↓ subst. G for A in the 5th pos.
ATGCCAT
↓ substitute C for G in the 3rd pos.
ATCCGAT

no. edit dist

(TGCATAT, ATCCGAT) = 4

∴ edit dist. = 4

→ 5 Edit operation:

$TG C A T A T \rightarrow TG C - A T A T$ (I → Insert)
 $A T C C G A T - -$ (D → Delete, S → Subst.)

→ 4 Edit operation:

$T G C A T A T$
 $A T C C G - A T$

⇒ Strong matching problem with insertion of gaps (indels) in any of the strings (with the constraint that simultaneous insertions at any position are not allowed) and also substitution of symbols. (a cost factors are generally associated with all these operations).

Given m -string S_m & n -string S_n , $\therefore$ ~~max~~ maximum possible edit operation $\Rightarrow \frac{m+n}{2}$ (only insertion & deletion) n gaps

Minimum edit operation $= (n-m)$, $n > m$. (only deletion & subst.)

$\therefore (n-m) < \text{edit-distance}(S_m, S_n) < \frac{m+n}{2}$
 However it is only D & I
 fighter board $\rightarrow$ max. (m, n)

⊕ Modeling edit operations:

- Two pointers movement (for each string one pointer)
 - ~~start~~ Initialize both the pointers in front of starting alphabet
 - move either one of them from its current position to the next position (as gap or indel is introduced in front of the static pointer.) A gap in the first reference string is insertion else (target string) it is deletion
 - move both the pointers (substitution / matching operation)
 - Stop when both the pointers arrived at the last alphabet.

Increased posn.
e.g. Ref: ~~P~~ T G C A T A T

Target: ~~P~~ A T C C G A T

Step 1. move P_t :

P_r ↑ — T G C A T A T
A T C C G A T

Edit
operation
I

Step 2. move P_r & P_t :

P_r ↓ — T G C A T A T
A T C C G A T
 P_t ↑

M

Step 3. move P_r & P_t :

— T G C A T A T
A T C C G A T
 P_r ↓ ↓ ↓
 P_t ↑

S,
M,
S,

Step 4. — do —
Step 5. — do —

Step 6. move P_r

— T G C A T A T
A T C C G — A T
 P_r ↓
 P_t ↑

D

Step 7 & 8. move P_r & P_t :

— T G C A T A T
A T C C G — A T
 P_r ↓
 P_t ↑

M, M.

∴ Total no. of edit operations $I + 4M + 2S + D$

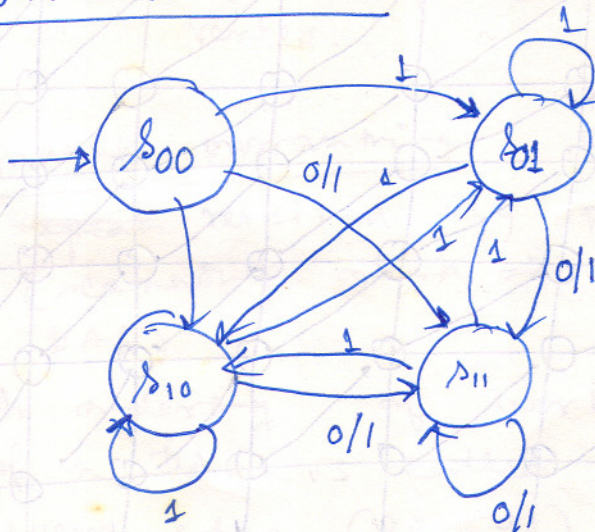
$\Rightarrow 4, M=0.$

N.B. (i) Movement is allowed from one alphabet to the next (jumping index)

(i) No. of $\downarrow p_r$ movements = n (length of ref. string)
 No. of $\uparrow p_t$ " = m (" of target ")

(ii) 00 $\Rightarrow$ No movement of $\downarrow p_r$ & $\uparrow p_t$ (only at initial state)
 01 $\Rightarrow$ Movement of $\downarrow p_r$
 10 $\Rightarrow$ " " $\uparrow p_t$
 11 $\Rightarrow$ " " $\downarrow p_r$ & $\uparrow p_t$

A State transition model:



0/1 $\rightarrow$ Cost
 for matching 0
 else 1.

stop after when
 $\sum 1 = m+n$

~~$n_0 + m$~~

Let n_{ij} = no. of times s_{ij} occurred.

$$\therefore n_{01} + n_{10} + 2 \cdot n_{11} = m+n; \quad n_{01} + n_{11} = n;$$

$$n_{10} + n_{11} = m;$$

(no. of integer solutions of the above eqn.)

gives the no. of valid sequences.

$\Rightarrow$ constraints:

$$n_{01} + n_{11} = n;$$

$$n_{10} + n_{11} = m;$$

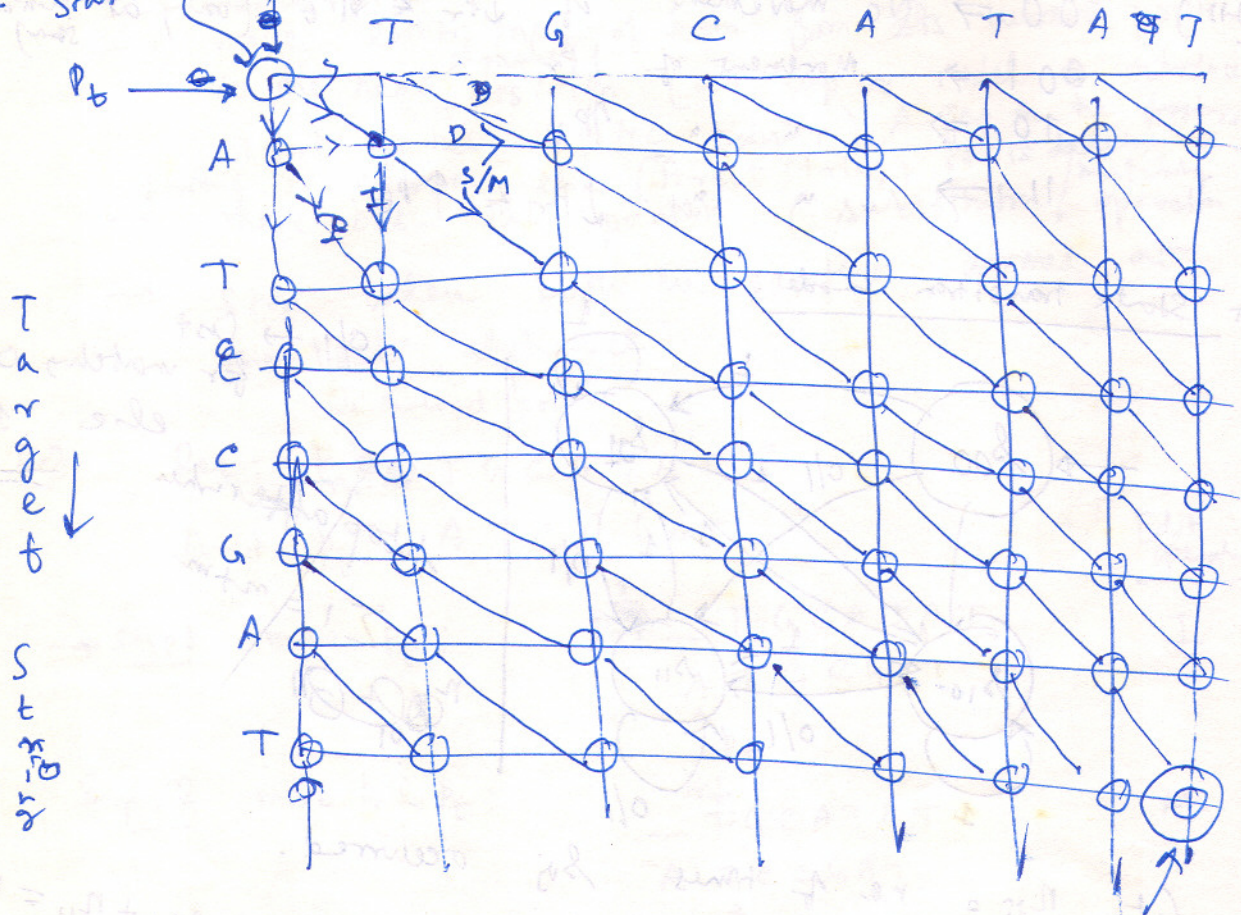
$$n_{01}, n_{10}, n_{11} \geq 0;$$

Content free
 transitions
 are not
 able to
 capture the
 accept
 or valid
 string.

It is quite difficult to ~~consider~~ ^{enumerate} all possible
 valid sequences & get the minimum out
 of it. States should be defined from the
 context.

Context Sensitive PSM: (Position of $\downarrow$ Pr & $\uparrow$ Pr define a State).
 Ref. sm

e.g. Starts state, Pr



Transitions are possible only to right ^{stop state} neighbor, bottom / diagonal neighbor for D/I/S/M operations.

⇒ Computation of edit distance is ^{equivalent to} the shortest path from the starting state to the stop state following the neighborhood definition.

⇒ Directed Acyclic Graph (DAG).

- Alg.:
- ① Start from the starting vertex and initialise distance to its neighbours vertices.
 - ① For every vertex U ,
 - if its predecessors are visited, mark it ~~computed~~ visited & compute its dist. from the source as $\min \{ \text{dist. (Predecessor)} + \text{cost} \}$
 - between $(U, \text{Pred.})$ & also store the predecessor info. ^{iteratively}
 - ② Perform step 1 till you reach destination vertex.

Dynamic programming: Ensure no vertex is ~~revisited~~ revisited by proper traversal order (scanning order $\Rightarrow$ topological ^{sorted} ordering for a DAG)

An ordering of vertices $v_1, v_2, \dots, v_n$ of a DAG is called topological if every edge (v_i, v_j) of the DAG connects a vertex with a smaller index to a vertex with a larger index, i.e., $i < j$.

$\rightarrow$ In this case row-wise scanning / column-wise scanning / zigzag scanning should be ok. $\Rightarrow O(m.n)$

$\rightarrow$ For every vertex store also the predecessor from which it is the shortest path from the source. This helps in retracing the path back.

$\rightarrow$ If weights are ^{related} inversely proportional to dist., then ~~longest~~ ^{maximum} path will provide the solution (maximisation problem). $\Rightarrow$ Used for sequence alignment.

* Longest Common Subsequence problem:

A subsequence of a string v is simply an (ordered) sequence of alphabets (not necessarily consecutive) from v .

e.g.

$v = \text{ATTGCTA}$

$A, G, C, A, A, T, T, A \in \text{set of subsequences of } v$
but $T, G, T, T, T, C, G \notin$

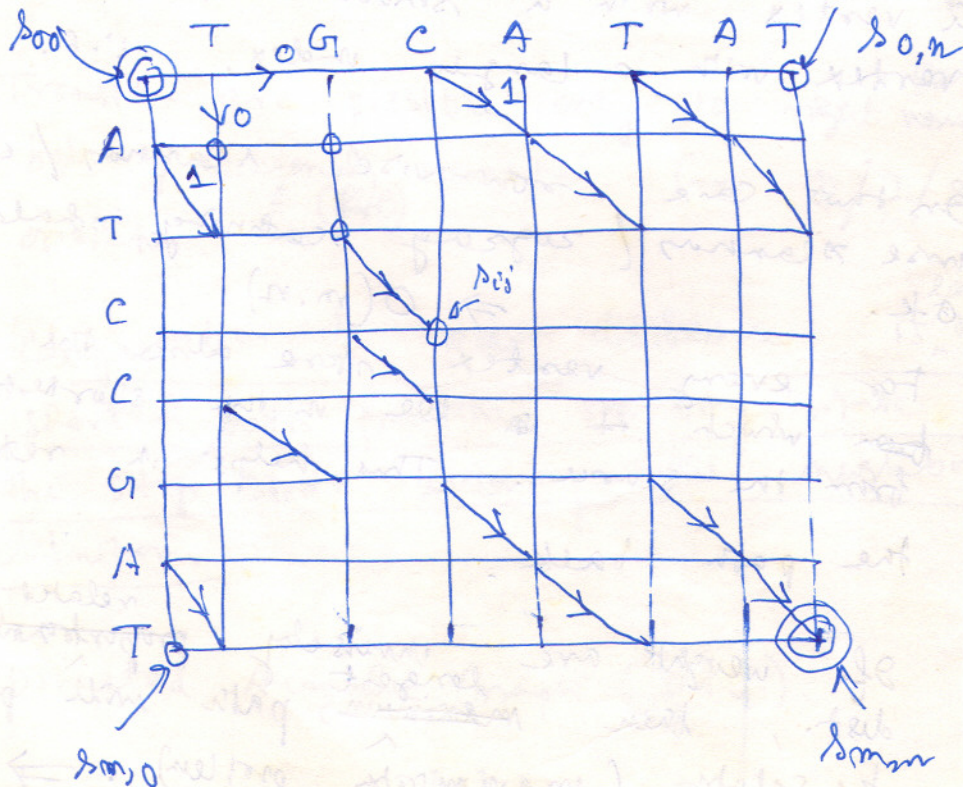
Problem Statement:

Let $v = v_1 v_2 \dots v_n$ & $w = w_1 w_2 \dots w_m$ be two strings. Find the longest common subsequence of v & w .

s.t. $v_{i_t} = w_{j_t}$

& $0 \leq i_1 < i_2 < \dots < i_n, 1 \leq j_1 < j_2 < \dots < j_m$.

$\Rightarrow$ No mutation or substitution allowed.



$$\begin{cases} \delta_{i,0} = \delta_{0,j} = 0, & 1 \leq i \leq m, \quad 1 \leq j \leq n \\ \delta_{i,j} = \max \begin{cases} \delta_{i-1,j} + 0 \\ \delta_{i,j-1} + 0 \\ \delta_{i-1,j-1} + 1, & \text{if } v_i = w_j \end{cases} \end{cases}$$

* B. (c) enforce back tracking pointers $\leftarrow, \uparrow, \rightarrow$ at every vertex.

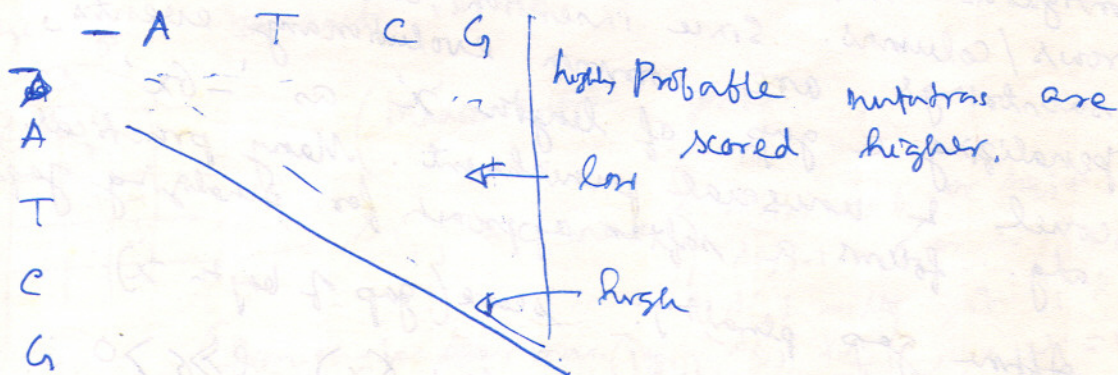
(b) Computation of edit distance:

$$d_{i,0} = i; \quad d_{0,j} = j; \quad 1 \leq i \leq m, \quad 1 \leq j \leq n;$$

$$d_{i,j} = \min \begin{cases} d_{i-1,j} + 1 \\ d_{i,j-1} + 1 \\ d_{i-1,j-1} + 0 & \text{if } v_i = w_j \end{cases}$$

* Global Sequence Alignment: Needleman & Wunsch (1970)

(The use of context sensitive FSM with all transitions).
scoring matrix ($\delta(\quad)$)



For protein comparison 21x21 scoring matrix.

PAM (Point accepted Mutation) $\rightarrow$

BLOSUM (Block substitution) $\rightarrow$

$$S_{0,j} = \max \begin{cases} S_{i-1,j} + \delta(v_i, -) \\ S_{i,j-1} + \delta(-, w_j) \\ S_{i-1,j-1} + \delta(v_i, w_j) \end{cases}$$

if, $\delta(v_i, -) = \delta(-, w_j) = -\sigma$ [insertion & deletion]
 $\delta(v_i, w_j) = -\mu$ (Mismatch), for $v_i \neq w_j$
 $= 1$ (match).

$$S_{i,j} = \max \begin{cases} S_{i-1,j} - \sigma \\ S_{i,j-1} - \sigma \\ S_{i-1,j-1} - \mu, \text{ if } v_i \neq w_j \\ S_{i-1,j-1} + 1, \text{ if } v_i = w_j \end{cases}$$

Resulting score = no. of matches $- \mu \times$ no. of mismatches
 $- \sigma \times$ no. of insertion & deletion.

Alignment with 'Gap' penalties:

A gap in an alignment is defined as a contiguous sequence of spaces (indels) in one of the rows/columns. Since insertion and deletions of substrings are common evolutionary events, penalizing a gap of length 'x' as $-\sigma x$ is cruel & unusual punishment. Many practical alg. follows a softer approach for penalizing gaps.

Affine gap penalty: score(gap of length x)

$$= -(p + \sigma x), \text{ where } p \geq 0, \sigma > 0$$

$$\downarrow$$

$$S_{i,j} = \max \begin{cases} S_{i,j-1} - \sigma \\ S_{i-1,j} - (p + \sigma) \end{cases}$$

$$\uparrow$$

$$S_{i,j} = \max \begin{cases} S_{i,j-1} - \sigma \\ S_{i-1,j} - (p + \sigma) \end{cases}$$

$$\delta_{i,j} = \max \begin{cases} \delta_{i-1,j-1} + \delta(v_i, w_j) \\ \delta_{i,j} \\ \delta_{i,j} \end{cases}$$

N.B.

[For every (i,j) the vertex maintain $\delta_{i,j}$, $\delta_{i,j}$ & $\delta_{i,j}$ separately].

* Local Sequence Alignment problem: Smith-Waterman (1981)

Problem statement: Given strings v & w and a scoring matrix δ , compute substrings of v & w whose global alignment, as defined by δ , is maximal among all global alignments of all substrings of v & w .

Model: Additional edges from the starting state (vertex) to each state (vertex) with "Zero" cost & also edges from the each vertex to the destination vertex.



(Maximisation problem)

$$\therefore \delta_{i,j} = \max \begin{cases} 0 \\ \delta_{i-1,j} - \delta \text{ or } \delta_{i,j} + \delta(v_i, -) \\ \delta_{i,j-1} - \delta \text{ or } \delta_{i,j-1} + \delta(-, w_j) \\ \delta_{i,j} - \delta, \text{ if } v_i \neq w_j \\ \delta_{i,j} + 1, \text{ if } v_i = w_j \end{cases}$$

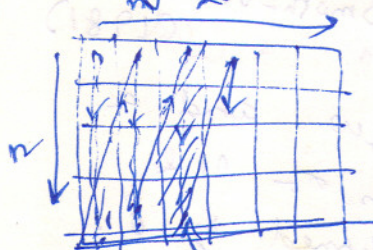
(Smith-Waterman (1980)).

$\delta(v_i, w_j)$

* Space Efficient Alignment Algo.:

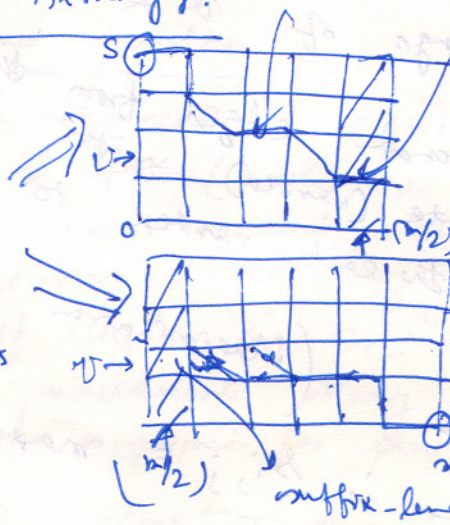
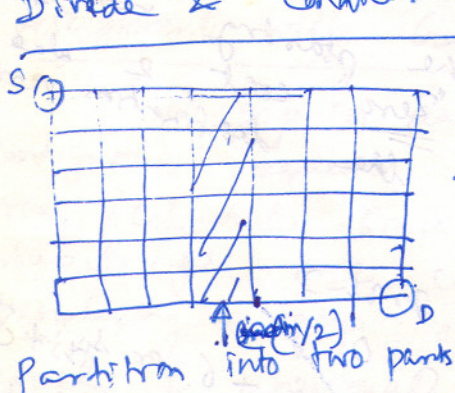
Divide & Conquer Strategy:

→ Usual Dynamic Algo.: Time complexity: $O(nm)$ ~~no. of~~ edges.
 Storage: $O(nm)$ ~~no. of~~ vertices.
 [for storing back-back pointers].
 However, if only 'score' is required (not the aligned sequence), $O(n \text{ or } m)$ is sufficient. Store only previous row or column during scanning.



Time complexity: $O(nm)$ } When only maximal score is required (or edit distance required).
 Space: $O(n)$ }
 remember only these ~~state~~ vertices for computing scores at next column.

→ Divide & Conquer Strategy:



Prefix-length $(i, m/2)$

Compute the maximal score at every vertex of the shaded cell (from S)

- same - for D.

suffix-length $(i, m/2)$.

$\therefore$ maximal length $(S, D) \geq \max \left(\text{Pref-length}(i, m/2) + \text{suffix-length}(i, m/2) \right)$
 & the vertex $(i, m/2)$ will be on the path for which it is max.

i.e.

$i = \text{argmax} (\dots)$

Computational requirement: $O(n \cdot m/2) + O(n \cdot m/2) \approx O(nm)$
 Storage: $O(n) + O(n)$.

6. Four Russian Speed Up: (Proposed by Vladimir Arlazarov, Effim Dinic, Mikhail Kronrod & Igor Paradzov in 1970, applied to sequence comparison by William Masek & Michael Paterson in 1980).

Store all possible combinations of t length substrings: (Storage space: $4^t \times 4^t \Rightarrow$ for nucleotides)

$$\text{Let, } t = \frac{\log n}{4}, \quad \text{}$$

$$\therefore n = 2^{4t} = 4^{2t}.$$

$$\therefore 4^t \times 4^t = n^{\frac{1}{2}} \times n^{\frac{1}{2}} = n.$$

$\therefore$ Storage space requirement: $O(n)$ for all optimal paths.

This makes computation of $\beta_{0,j}$ for accession t string $\Rightarrow O(\frac{n}{t}) \Rightarrow O(\log n)$

$$\therefore \text{Computation of all } \beta_{0,j} \Rightarrow O\left(\frac{n}{t} \times \frac{n}{t}\right) \times \log n$$

$$\Rightarrow O\left(\frac{n^2}{t^2}\right) \times \log n.$$

$$\Rightarrow O\left(\frac{n^2 (\log n)}{(\log n)^2}\right)$$

$$\Rightarrow O\left(\frac{n^2}{\log n}\right).$$

$\Rightarrow$ Discuss Doton chaining problem

(i) Splicing based Alignment.

* Heuristic Similarity Search Techniques:

→ Dynamic Algorithm: $O(n^2)$ not suitable for searching through large genome database (10^{10} nucleotides) & query string ($\sim 10^3$)

→ Starting in the early 1990's biologists had no choice but to use fast heuristics as an alternative to quadratic sequence algorithm.

→ D. Lipman & W.R. Pearson (1985) → FASTA

S. Altschul, W. Gish, W. Miller, E. Myers and J. Lipman

→ (1990) → Basic Local Alignment Search Tool (BLAST)

→ Modified in 1997.

• l-mer filtration technique for approximate pattern matching (with k -mismatches)

If an n -letter substring of a query approximately matches an n -letter substring of the text, then the two substrings share at least one l -mer for a sufficiently large value of l .

Th.: If the strings $x_1, \dots, x_n$ & $y_1, \dots, y_n$ match with at most k mismatches, then they share an l -mer for $l = \lfloor \frac{n}{k+1} \rfloor$, i.e. $x_{i+1} \dots x_{i+l} = y_{i+1} \dots y_{i+l}$ for some $1 \leq i \leq n-l+1$

Proof: Partition into $(k+1)$ groups, $\{1 \dots l\}, \{l+1, \dots, 2l\}, \dots$

— where $l = \lfloor \frac{n}{k+1} \rfloor$;

As there are k mismatches, one of this set will not have any mismatch. □

This theorem motivates the following l -mer filtration algorithm for query matching with k -mismatches:

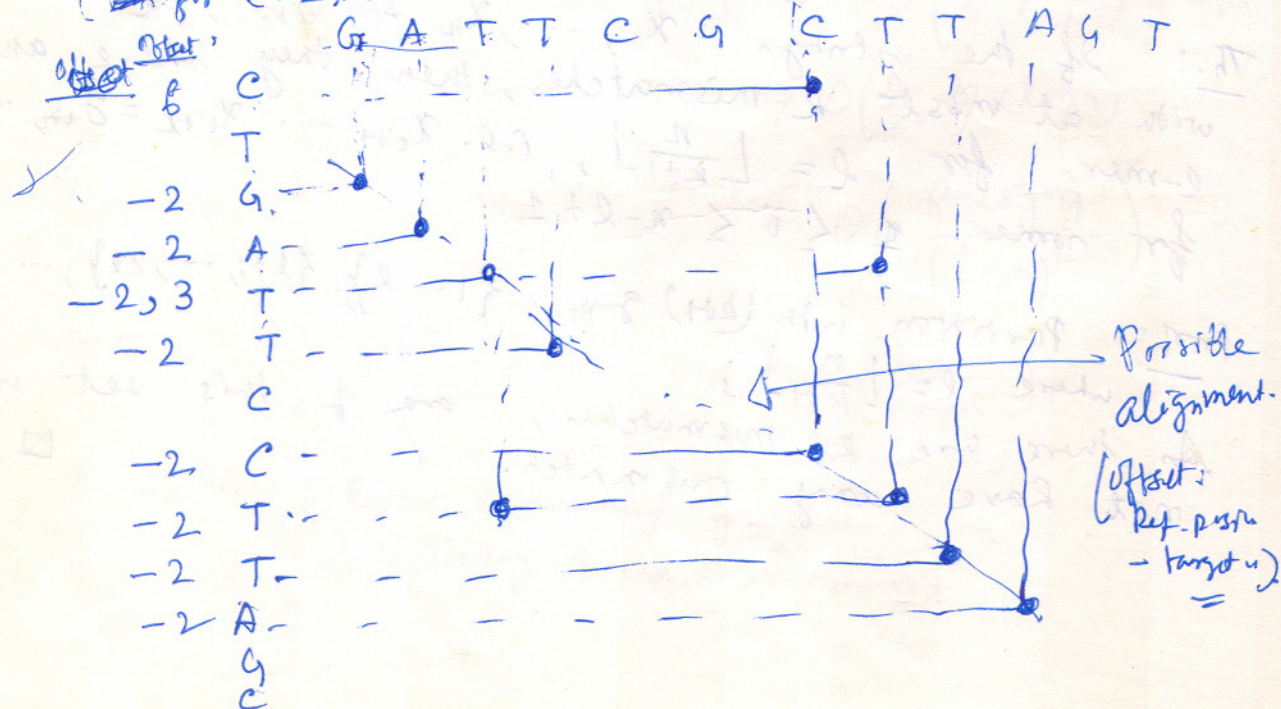
- Potential match detection: Find all matches of l -mers in both the query & the text for $l = \lfloor \frac{n}{k+1} \rfloor$. [Use of hashing or suffix tree]
- Potential match verification: Verify each potential match by extending it to the left & to the right until the first $k+1$ mismatches are found or the beginning/end of the query/text is found.

• FASTA:

→ FASTA search begins by breaking the search sequence into words (for nucleotides words of length 4 to 6, for proteins 1 to 2).

→ Find the positions of l -mer matches & represent in the form of dot matrices.

(In for $l=2$):



→ From this dot matrix, it selects some diagonal with a high concentration of 1's and groups these runs of ^{diagonal} 1's into longer runs. [~~By~~ ^{eg}, a local sequence alignment on these segments by dynamic programming).

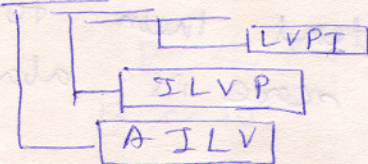
→ 'dots' are noted as offsets from the ref. positions. ^{Adjacent} Alphabets with similar offsets are grouped together.

• BLAST:

Input sequence: AILVPTV

1. Break the query sequence into words.

AILVPTVIGCT



} select those words which are biologically more relevant e.g. sequences with frequently observed amino acids are discarded.

2. Search for exact word matches (also called high-scoring pairs or HSPs) in the database sequence.

AILVPTV →
MVQGWALYDFLKCR AILVGT V I AML...

(in both directions)

3. Extend the match - until the local alignment score falls below a fixed threshold (the most recent version of BLAST allows gaps in the extended match).

→ For any two l -mers $x_1, \dots, x_l$ & $y_1, \dots, y_l$,
 BLAST defines the segment score as
 $\sum_{i=1}^l \delta(x_i, y_i)$, $\delta(x, y)$ is the similarity score
 between x & y .

→ A segment pair is locally maximal if
 its score cannot be improved by shifting
 it to any direction. BLAST attempts to
 find all locally maximal segment pairs in
 the query & database sequences with scores
 above some set threshold. If the threshold
 is high enough, then the set of all l -mers
 that have score above a threshold is not
too large. This becomes a multiple-pattern
 matching problem (use of keyword tree / suffix tree).

→ After the potential matches are located,
 BLAST attempts to extend them to see
 whether the resulting score is above the
 threshold.

→ The choice of the threshold in BLAST
 is guided by the Altschul-Dembo-Karlin
 statistics.

$$E(0) = \text{Number of matches with scores above } 0$$

$$= K \cdot m \cdot n \cdot e^{-\lambda \theta}$$

$m, n \Rightarrow$ length of two strings,
 $K \Rightarrow$ constant
 $\lambda \Rightarrow$ Parameter

λ is a free root of the eqn.:

$$\sum_{x, y \in A} p_x p_y e^{\lambda \delta(x, y)} = 1$$

$p_x, p_y \Rightarrow$ frequencies of x & y .

The probability that there is a match of score greater than 0 betw. two "random" sequences of length n & m is computed as $1 - e^{-E(0)}$. (Is it $1 - e^{-E(0)}$?)

* Alignment Scores and Statistical Significance of Database Searches:

E-score(S): expected no. of sequences of score $\geq S$, would have been found randomly.

P-score(S): prob. that one or more sequences of score $\geq S$ would have been found randomly.

Low values of E & P indicate the search result was unlikely to have been obtained by random chance & thus is likely to bear evolutionary relationship to the query sequence.

e.g. $E < 10^{-3} \Rightarrow$ indicative of stat. signif.

$E \sim 10^{-50} \Rightarrow$ strong indication.

* Multiple Sequence Alignment

→ Extend pairwise sequence alignment for k multiple sequences with dynamic programming in a k -dimensional space $\Rightarrow O((2n)^k)$

→ Heuristic method: (i) $\binom{k}{2}$ alignment & combine them pair
or (ii) get the best alignment & merge them into a new string (once a gap is always a gap).
Progressively, merge the best possible alignment with the combined string.
e.g. CLUSTAL (Higgins, Thompson & Gubson (1996))

• CLUSTAL (Higgins & Sharp, 1988) :

→ Constructs a phylogenetic tree to determine the degrees of similarity among the sequences being aligned.

→ Using this tree as a guide, closely related sequences are aligned two at a time using DP of the pairwise alignments.

→ Sensitive to initial choice of pairs.

→ CLUSTALW (1996) uses sequence weighting, position-specific gap penalties & weight matrix choice based on how the sequences are divergent.