

Sequence Alignment:-

Is a way of arranging of two or more sequences of DNA, RNA or protein to identify and highlight regions of similarity. The similarity may indicate the functional, structural and evolutionary significance of the sequence. It's made between a known sequence and unknown sequence or between two unknown sequences.

i - the known sequence is called reference sequence.

ii - the unknown sequence is called Query sequence.

TYPES OF ALIGNMENT

According to the number of sequence.

i - pair-wise alignment (Global or local sequence alignment)

ii - multiple alignment.

PAIR-WISE Sequence Alignment.

- it only compares two sequence at a time
- it is concerned with finding the best-matching of protein amino acid or DNA (nucleic acid) sequences.

- typically, the purpose of this is to find homologues (relative) of a gene in a database of knowns.

- this information is useful for answering a variety of biological questions.

i - the study of molecular evolution.

- It consists of a series of paired bases, one base from each sequence

TYPES OF PAIRS:

- i - matches: the same nucleotide appears in both sequences.
- ii - mismatches: different nucleotides are found in the two sequences, it's also alignment correspond to mutation.
- iii - Gaps: A base in one sequence and a null base in the other. Gaps correspond to insertions or deletions.

Difference between pairwise and multiple sequence Alignment:

- Pair-wise sequence Alignment:** - Algorithm used are -
- Needleman-Wunsh algorithm & Smith-Waterman algorithm.
 - Blast (Basic local alignment search tool).
 - it determine similar regions between two nucleotide or protein sequences; Optimality is based on score.
 - The BLAST algorithm is much more efficient, much faster ease of use, statistical rigor than other approaches.

Multiple sequence Alignment: (MSA): is a sequence alignment of three or more biological sequences, generally protein, DNA or RNA. It try to align all of the sequences in a specified set. The most popular multiple alignment tools is; ClustalW, MUSCLE, PRABCONS, MAFFT, and T-coffee.

TYPE OF Sequence Align

According to sequence coverage:

- i- Global sequence alignment
- ii- Local sequence alignment.

Global Sequence Alignment: Is a matching the residues of two sequences across their entire length. That is alignment is carried out from beginning to end of both sequences to find the best possible alignment across the entire length between the two sequences.

Local Sequence Alignment: Is a matching of two sequence from regions which have more similarity with each other, it does not assume that the two sequence in question have similarity over the entire length. It only find local regions with the highest level of similarity between the two sequences and aligns these regions without regard for the alignment of the rest of the sequence regions.

DATABASE SIMILARITY SEARCHING.

A main application of pair wise alignment is retrieving biological sequences in databases based on similarity; this process involves submission of a query sequence and performing a pair wise comparison of the query with all individual sequence in a database.

(Mycobacterium tuberculosis DNA, complete genome strain H37Rv)

There are two major algorithms for performing database searches: BLAST and FASTA.

BLAST (Basic Local Alignment Search Tool) is the most common local tool developed by Altschul et al. in 1990. It is an algorithm for comparing primary biological sequence information, such as the amino-acid sequences of different proteins or the nucleotides of DNA sequences. A BLAST search enables a researcher to compare a query sequence with a library or database of sequences, and identify library sequences that resemble the query sequence above a certain threshold.

Example the question that BLAST ask!

"Which bacterial species have a protein that is related in lineage to a certain protein with known amino-acid sequence"?

The BLAST algorithm and the computer program that implements it were developed by Stephen Altschul, Warren Gish and David Lipman at the ~~USA~~ US National Center for Biotechnology Information (NCBI)

FASAT FORMAT

IS The standard format in the field of bioinformatics to represent either nucleotide sequences or peptide sequences. this format is single-letter code and it allows sequence names and description at the beginning followed by sequence data in multiple lines. The length of the each chunk (line) of the sequence must not exceeds 80 characters, sequence identifiers are defined by a standard called NCBI