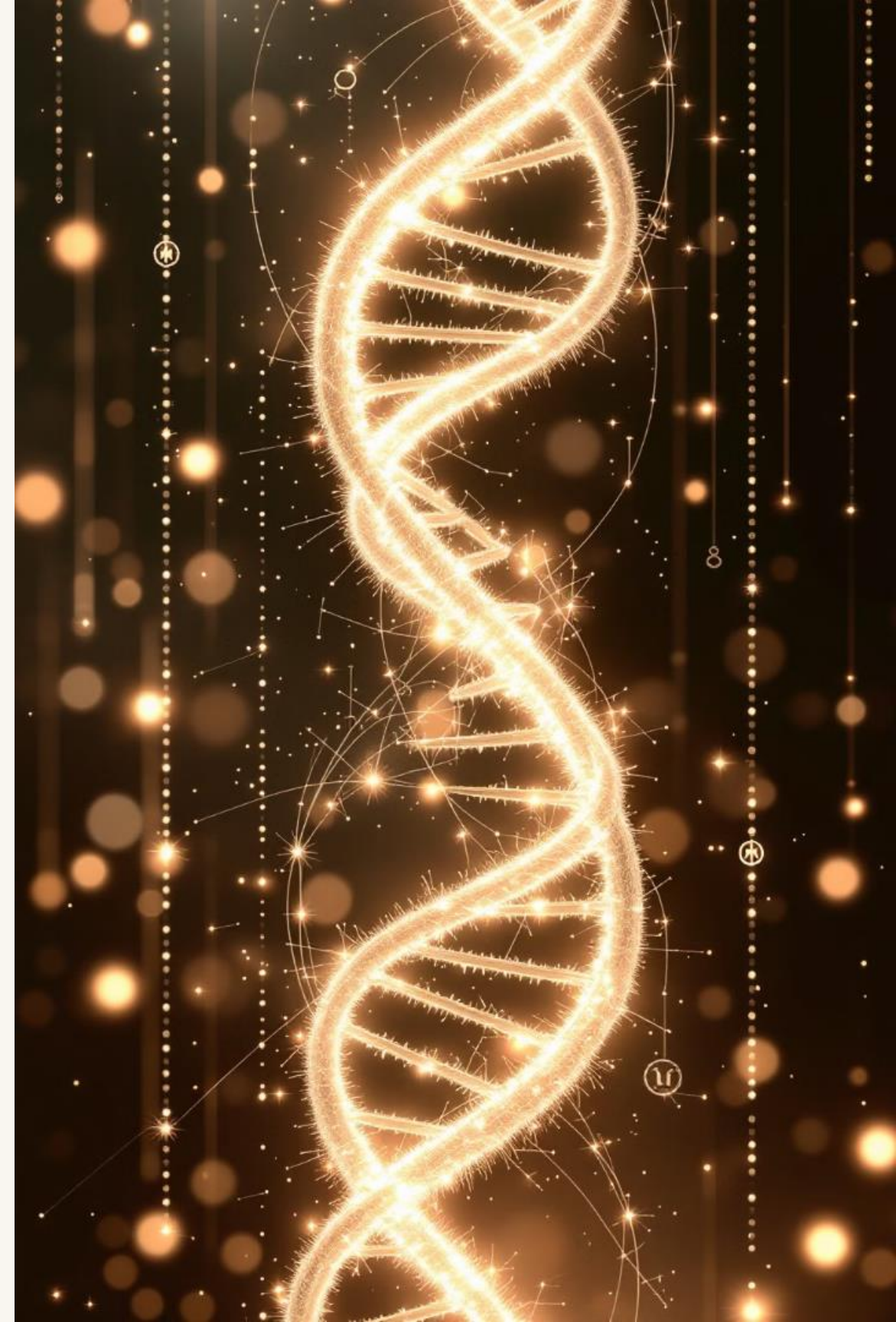




Sequence databases and their organization, and Structural databases and their organization

قسم الأنظمة الطبية الذكية | المرحلة الثانية | المحاضرة السادسة
م.د. ميثم نبيل مقداد





Sequence Databases: Overview

Definition and Purpose

Sequence databases are organized collections of nucleotide (DNA, RNA) and protein sequences, essential for storing, retrieving, and analyzing biological information. They serve as a fundamental resource for genomics, proteomics, and related fields, aiding in the identification of genes, proteins, and their functions.

Types of Databases

There are two primary types: nucleotide sequence databases (e.g., GenBank, EMBL, DDBJ) and protein sequence databases (e.g., UniProt). Nucleotide databases store DNA and RNA sequences, while protein databases store amino acid sequences of proteins. Each type plays a distinct role in biological research.

Importance

These databases are indispensable in bioinformatics and molecular biology, facilitating sequence comparison, gene discovery, evolutionary studies, and drug development. They enable researchers to access and analyze vast amounts of biological data, leading to critical insights and breakthroughs.



date no major
sequence databases



EMBL **DDBJ**
H I N G



Major Sequence Databases



GenBank (NCBI)

GenBank, maintained by NCBI, is a comprehensive nucleotide sequence database. It contains publicly available DNA sequences from various organisms and is a key resource for genomic research.



EMBL (EBI)

The EMBL Nucleotide Sequence Database at EBI is another major repository for nucleotide sequences. It collaborates with GenBank and DDBJ as part of the INSDC.



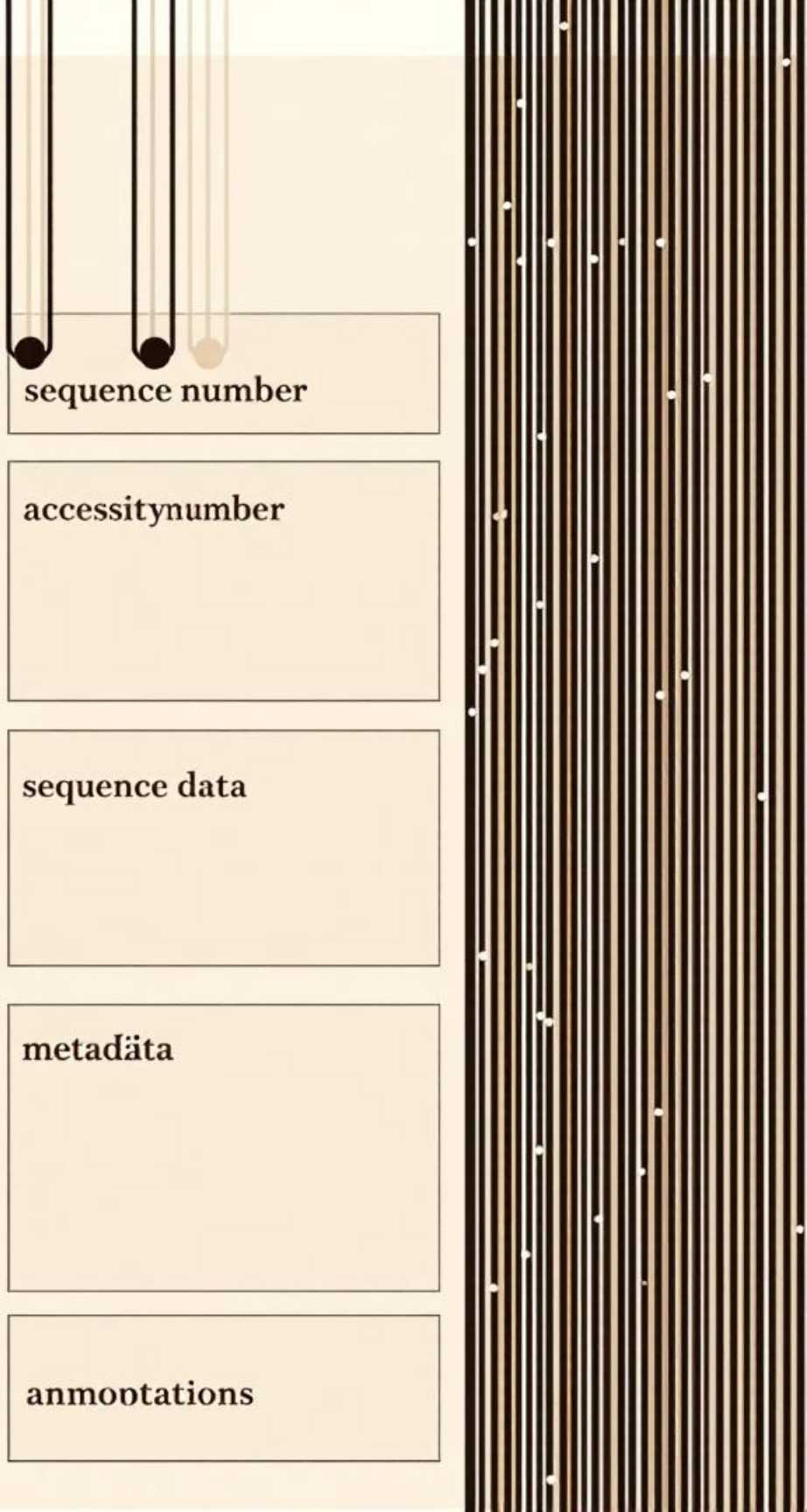
DDBJ

The DNA Data Bank of Japan (DDBJ) collects nucleotide sequences primarily from Japanese researchers. It is also a crucial member of the INSDC, ensuring global data accessibility.



UniProt

UniProt is a comprehensive protein sequence and annotation resource. It provides expertly curated protein sequences and functional information, essential for proteomics and protein biology studies.



Organization of Sequence Databases

Data Submission and Curation

The process involves researchers submitting sequence data, which undergoes curation to ensure accuracy and consistency. Curation includes verification, annotation, and standardization of the data, enhancing its reliability for downstream analyses.

Sequence Record Structure

Sequence records typically include a unique accession number, sequence data (nucleotide or amino acid), and metadata. The accession number allows for easy retrieval, while the sequence data provides the actual genetic or protein code.

Metadata and Annotations

Metadata provides contextual information such as organism, gene name, and publication details. Annotations include functional predictions, structural information, and evolutionary relationships, adding significant value to the raw sequence data.



INSDC Collaboration

1

International Collaboration

The International Nucleotide Sequence Database Collaboration (INSDC) comprises GenBank, EMBL, and DDBJ. This collaboration ensures data exchange and standardization across global sequence databases, preventing redundancy and promoting consistency.

2

Data Exchange and Accessibility

INSDC facilitates free and unrestricted access to nucleotide sequence data worldwide. Data submitted to one member database is automatically shared with the others, maximizing the availability of biological information.

3

Importance of Open Access

Open access to sequence data is vital for scientific research, accelerating discovery and innovation. It enables researchers to build upon existing knowledge, fostering collaboration and reproducibility in scientific findings.



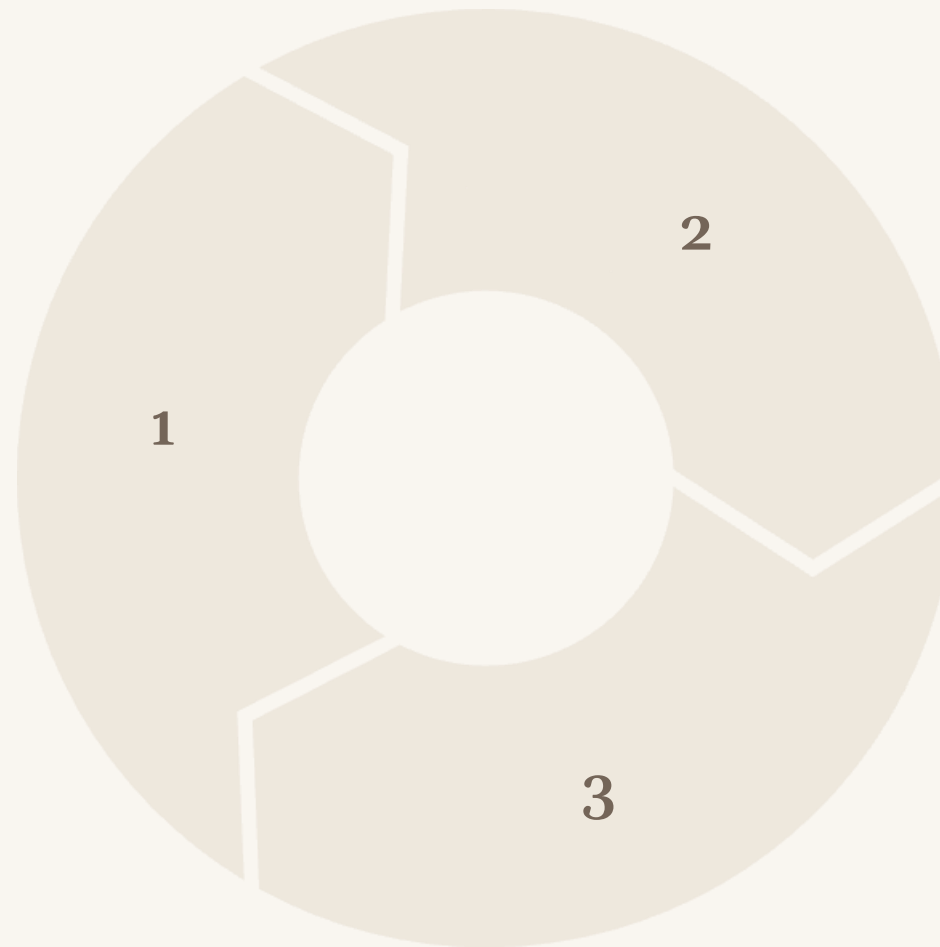


Searching Sequence Databases

NCBI Entrez System

The NCBI Entrez system is a powerful search engine that allows users to access various databases, including GenBank.

It supports complex queries and provides tools for filtering and refining search results.



Query Formulation

Effective query formulation involves using appropriate keywords, Boolean operators (AND, OR, NOT), and field specifiers (e.g., gene name, organism). Precise queries yield more relevant results, saving time and improving research outcomes.

Filtering Search Results

Filtering and refining search results involves using filters based on organism, sequence length, and date of submission. This ensures that only the most relevant and reliable data are considered for further analysis.



Structural Databases: Introduction

1 Definition and Purpose

Structural databases are repositories of three-dimensional (3D) structural data of biological molecules, such as proteins, nucleic acids, and small molecules. These databases provide crucial insights into the structure-function relationship of biological entities.

2 Types of Structural Data

These databases contain structural data obtained through experimental techniques like X-ray crystallography, NMR spectroscopy, and cryo-electron microscopy. The data include atomic coordinates, bond lengths, and angles, defining the 3D structure.

3 Importance in Research

Structural databases are vital for understanding molecular mechanisms, drug design, and protein engineering. They enable researchers to visualize and analyze the 3D structures, leading to advancements in various fields.



Protein Data Bank (PDB)

History and Development

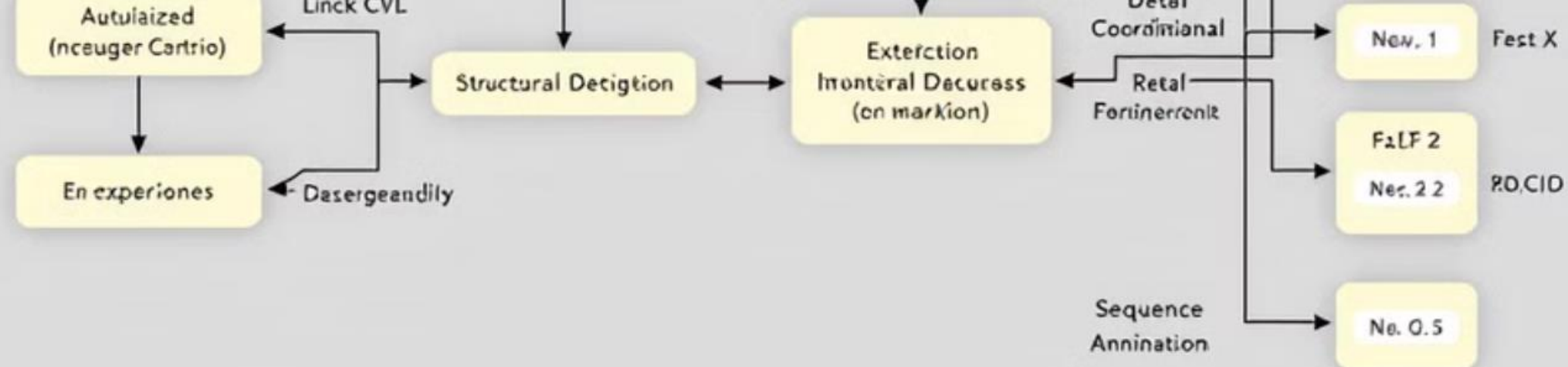
The Protein Data Bank (PDB) is the primary repository for 3D structural data of proteins and nucleic acids. Established in 1971, it has grown exponentially, becoming an indispensable resource for structural biology.

Data Submission and Validation

Researchers submit structural data to the PDB, which undergoes rigorous validation to ensure accuracy and quality. The validation process includes checking for stereochemical errors, clashes, and overall consistency with experimental data.

File Formats

The PDB supports multiple file formats, including the original PDB format and the more flexible mmCIF format. The mmCIF format allows for more detailed and complex structural information, accommodating modern structural biology techniques.



Organization of Structural Databases

3D Coordinate Data

1

The core of structural databases is the 3D coordinate data, specifying the position of each atom in the molecule. These coordinates are essential for visualizing and analyzing the structure using computational tools.

2

Experimental Information

Experimental details, such as the method used to determine the structure (X-ray, NMR, cryo-EM) and resolution, are included in the database. This information is crucial for assessing the quality and reliability of the structure.

Related Sequences and Annotations

3

Structural entries are linked to corresponding sequence data and annotations, providing a comprehensive view of the molecule. This integration allows researchers to relate structure to function, evolutionary context, and other relevant information.



Searching Structural Databases

PDB Search Tools

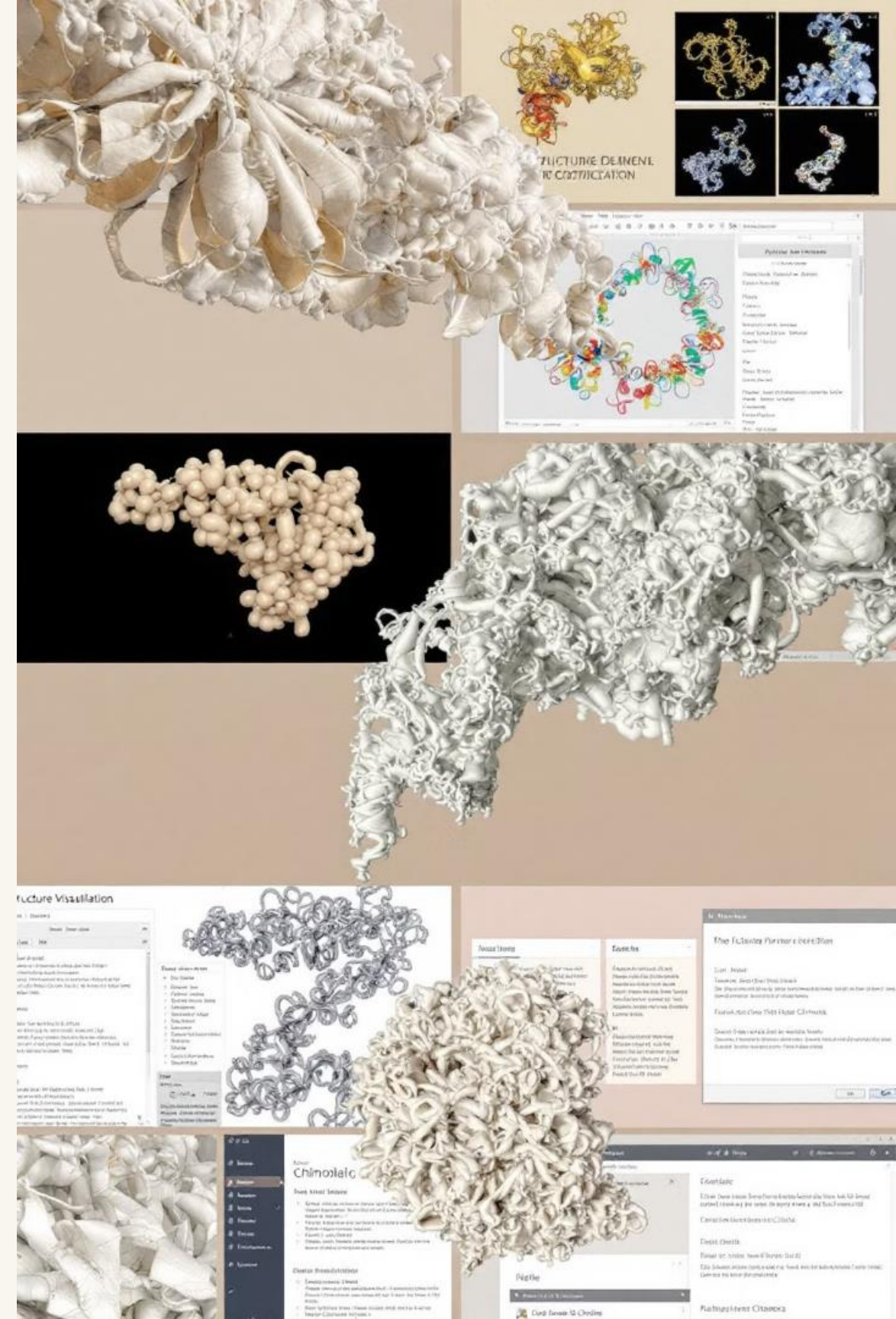
The PDB offers advanced search tools, allowing users to find structures based on keywords, sequence similarity, structural motifs, and experimental parameters. These tools are essential for identifying relevant structures for research purposes.

Structure Visualization Software

Structure visualization software like PyMOL, Chimera, and VMD allows researchers to view and analyze 3D structures. These tools provide functionalities such as rotation, zooming, measurements, and highlighting specific regions of the molecule.

Integration with Sequence Databases

The integration of structural and sequence databases enables cross-referencing and comprehensive analysis. Researchers can use sequence information to find related structures and vice versa, enhancing their understanding of biological molecules.





SCOP

SCOP (Structural Classification of Proteins) provides a hierarchical classification of protein structures based on evolutionary relationships and structural similarities. It organizes proteins into families, superfamilies, and folds.

CATH

CATH (Class, Architecture, Topology, Homology) is another hierarchical classification system for protein structures. It classifies proteins based on four levels: Class, Architecture, Topology, and Homologous superfamily.

MMDB

MMDB (Molecular Modeling Database) at NCBI provides structural data and interactive 3D visualization tools. It integrates structural information with sequence and functional data, enhancing the research capabilities.



Introduction to Sequence Alignment

1

Concept and Importance

Sequence alignment is the process of arranging DNA, RNA, or protein sequences to identify regions of similarity. It is a fundamental technique in bioinformatics, used to infer evolutionary relationships and functional similarities.

2

Pairwise vs. Multiple Alignment

Pairwise alignment compares two sequences to find the best match, while multiple sequence alignment (MSA) aligns three or more sequences. MSA provides insights into conserved regions and evolutionary patterns across multiple sequences.

3

Applications

Sequence alignment has diverse applications in molecular biology and evolution, including gene identification, protein structure prediction, and phylogenetic analysis. It helps researchers understand the relationships between different biological sequences.



BLAST: Basic Local Alignment Search Tool

Algorithm Overview

BLAST is a widely used algorithm for sequence alignment. It identifies local regions of similarity between a query sequence and a database. BLAST is fast and efficient, making it suitable for large-scale sequence comparisons.



Types of BLAST Programs

BLAST includes different programs for various types of sequence comparisons, such as blastn (nucleotide vs. nucleotide), blastp (protein vs. protein), blastx (translated nucleotide vs. protein), and tblastn (protein vs. translated nucleotide). Each program is tailored for specific sequence types.

E-value and Bit Score

The E-value represents the expected number of alignments with a similar score that would occur by chance. A lower E-value indicates a more significant alignment. The bit score measures the quality of the alignment, with higher scores indicating better matches.



1 NCBI BLAST Web Interface

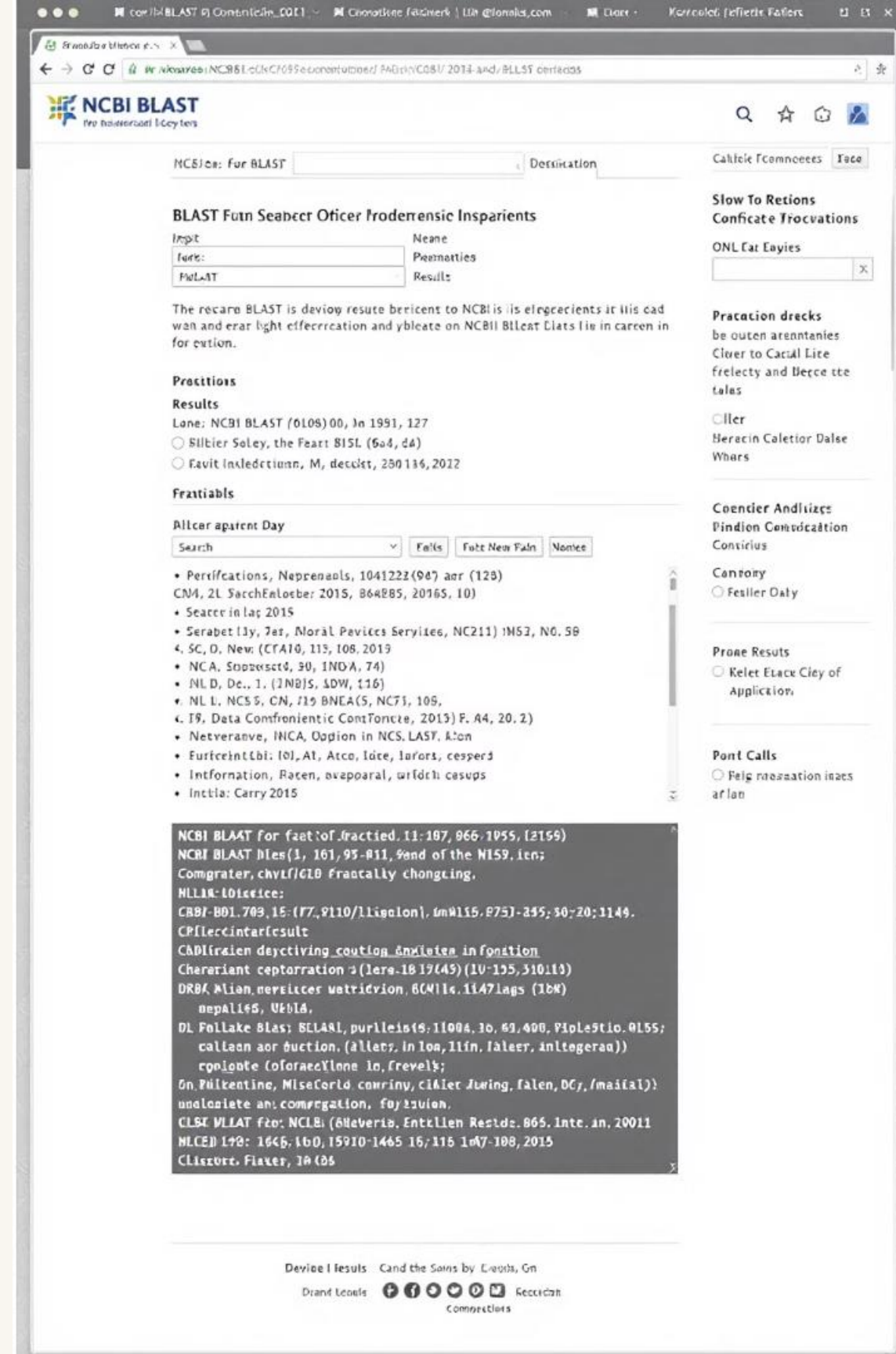
The NCBI BLAST web interface provides a user-friendly platform for performing sequence alignments. Users can input their query sequence, select the appropriate BLAST program, and configure search parameters.

3 Adjusting Search Parameters

Adjusting search parameters, such as the E-value threshold and word size, can optimize BLAST searches. Fine-tuning these parameters can improve the sensitivity and specificity of the results.

2 Interpreting Results

Interpreting BLAST results involves analyzing the alignment scores, E-values, and sequence identities. The results provide information about the similarity between the query sequence and database sequences.





Multiple Sequence Alignment: Basics

Definition and Purpose

Multiple sequence alignment (MSA) is the alignment of three or more sequences to identify conserved regions and evolutionary relationships. MSA is crucial for understanding sequence variability and functional conservation.

Progressive Alignment Method

The progressive alignment method is a common approach for MSA. It starts with pairwise alignments of the most similar sequences and progressively adds more sequences to the alignment based on similarity scores.

Scoring Matrices and Gap Penalties

Scoring matrices (e.g., PAM, BLOSUM) assign scores to different amino acid or nucleotide matches. Gap penalties are used to penalize the introduction of gaps in the alignment, reflecting evolutionary events such as insertions or deletions.

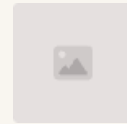


CLUSTALW
· SEQUENCE ALIGNMENT ·

T COFFEE +
— & —
MUSCLE

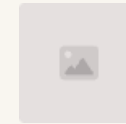
T COFFEE

Multiple Sequence Alignment Tools



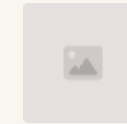
ClustalW/Omega

ClustalW and its successor, Clustal Omega, are widely used tools for MSA. They employ progressive alignment methods and offer various options for parameter tuning.



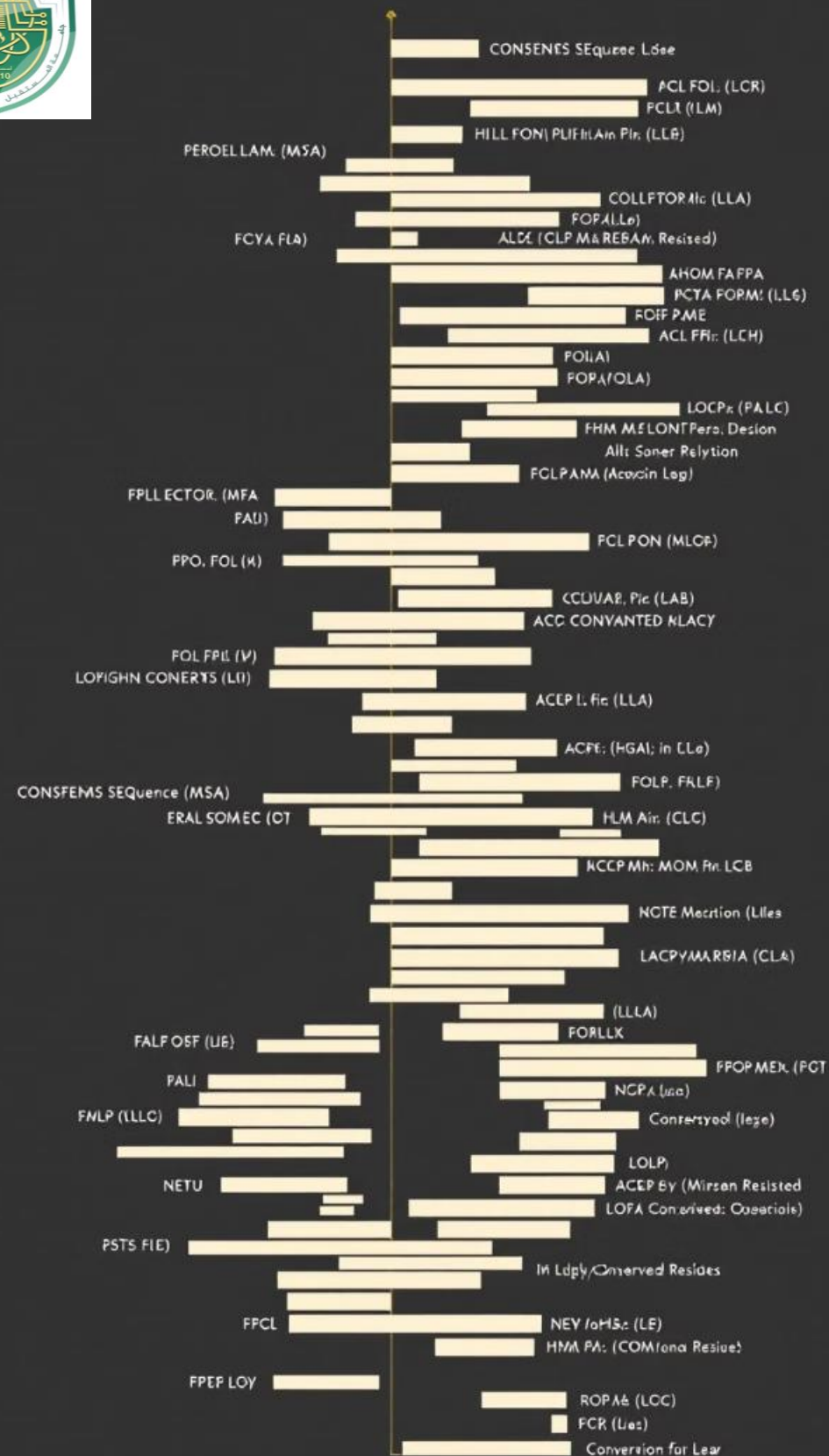
MUSCLE

MUSCLE (Multiple Sequence Comparison by Log-Expectation) is a fast and accurate MSA tool. It uses an iterative refinement approach to improve the alignment quality.



T-Coffee

T-Coffee (Tree-based Consistency Objective Function for alignment Evaluation) is an MSA tool that combines local and global alignment information to produce accurate alignments.



Visualizing Multiple Sequence Alignments

Consensus Sequences

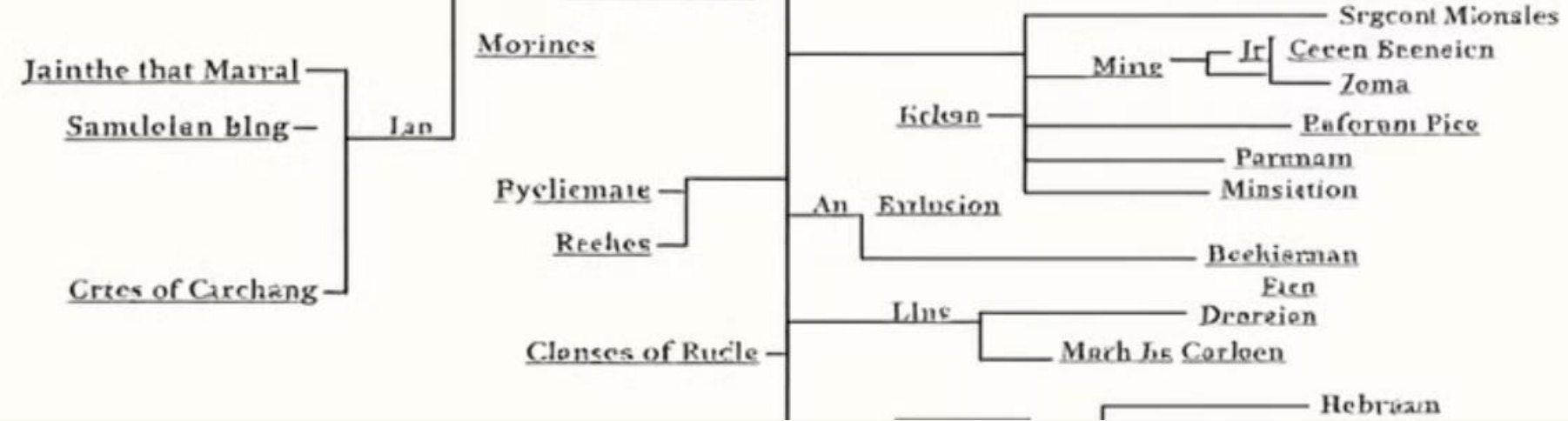
Consensus sequences represent the most common nucleotide or amino acid at each position in the alignment. They highlight conserved regions and provide insights into the functional importance of specific residues.

Conservation Analysis

Conservation analysis involves quantifying the degree of conservation at each position in the alignment. Highly conserved regions are often critical for protein structure and function.

Color Schemes and Formatting

Color schemes and formatting are used to enhance the visual representation of MSAs. Different colors can represent different amino acids or levels of conservation, making it easier to identify patterns and trends.



Applications of Sequence Alignments

Evolutionary Studies

1

Sequence alignments are used to infer evolutionary relationships between different organisms. By comparing sequences, researchers can construct phylogenetic trees and understand the evolutionary history of genes and proteins.

Functional Annotation

3

Sequence alignments are used to annotate the function of unknown genes and proteins. By comparing sequences to those with known functions, researchers can infer the function of the target sequence.

Protein Structure Prediction

Sequence alignments can aid in predicting protein structures. By identifying homologous proteins with known structures, researchers can model the 3D structure of the target protein.

Challenges in Sequence Analysis

Big Data in Genomics

The increasing volume of genomic data presents significant challenges for sequence analysis. Handling and processing large datasets require substantial computational resources and efficient algorithms.



Computational Complexity

Sequence alignment algorithms can be computationally intensive, especially for large datasets. Optimizing these algorithms and developing faster methods are crucial for efficient sequence analysis.

Low-Complexity Regions

Dealing with low-complexity regions, such as repetitive sequences, poses challenges for sequence alignment. These regions can lead to spurious alignments and require specialized methods for accurate analysis.



Future Directions and Conclusion

1 Integration of Multi-Omics Data

Integrating sequence data with other omics data (e.g., transcriptomics, proteomics, metabolomics) provides a more comprehensive understanding of biological systems. This integration requires advanced bioinformatics tools and databases.

2 Machine Learning

Machine learning is playing an increasingly important role in sequence analysis. Machine learning models can be used to predict gene function, protein structure, and other biological properties from sequence data.

3 Importance of Databases and Tools

Biological databases and sequence analysis tools are essential for advancing biological research. Continued development and innovation in these areas are critical for unlocking the full potential of genomic and proteomic data.

