

Getting To Know Your Protein

Comparative Protein Analysis: Part II. Protein Domain Identification & Classification

Robert Latek, PhD
Sr. Bioinformatics Scientist
Whitehead Institute for Biomedical Research

Comparative Protein Analysis

- **Part I. :**
 - **Phylogenetic Trees** and **Multiple Sequence Alignments** are important tools to understand global relationships between sequences.
 - <http://jura.wi.mit.edu/bio/education/bioinfo-mini/seq/> (AA Subst. Matrices)
- **Part II. :**
 - How do you identify sequence relationships that are restricted to localized regions?
 - Can you apply phylogenetic trees and MSAs to only sub-regions of sequences?
 - How do you apply what you know about a group of sequences to finding additional, related sequences?
 - What can the relationship between your sequences and previously discovered ones tell you about their function?
- Assigning sequences to **Protein Families**

WIBR Bioinformatics Courses, © Whitehead Institute, 2004

2

Syllabus

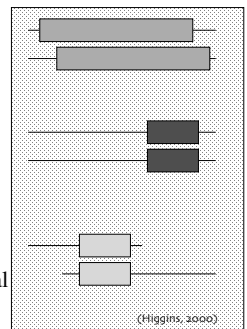
- **Protein Families**
 - Identifying Protein Domains
 - Family Databases & Searches
- **Searching for Family Members**
 - Pattern Searches
 - Patscan
 - Profile Searches
 - PSI-BLAST/HMMER2

WIBR Bioinformatics Courses, © Whitehead Institute, 2004

3

Proteins As Modules

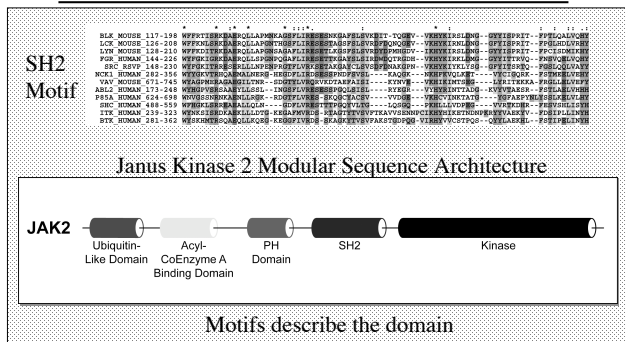
- Proteins are derived from a limited number of basic building blocks (**Domains**)
- Evolution has shuffled these modules giving rise to a diverse repertoire of protein sequences
- As a result, proteins can share a global or local relationship



WIBR Bioinformatics Courses, © Whitehead Institute, 2004

4

Protein Domains



WIBR Bioinformatics Courses, © Whitehead Institute, 2004

5

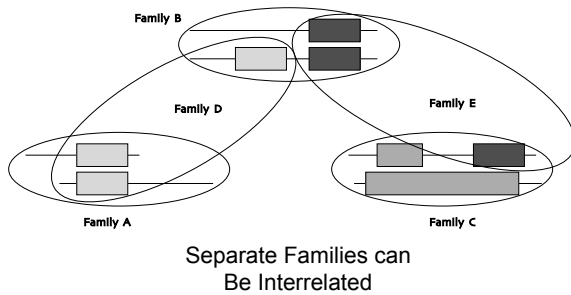
Protein Families

- **Protein Family** - a group of proteins that share a common function and/or structure, that are potentially derived from a common ancestor (set of homologous proteins)
- **Characterizing a Family** - Compare the sequence and structure patterns of the family members to reveal shared characteristics that potentially describe common biological properties
- **Motif/Domain** - sequence and/or structure patterns common to protein family members (a trait)

WIBR Bioinformatics Courses, © Whitehead Institute, 2004

6

Protein Families

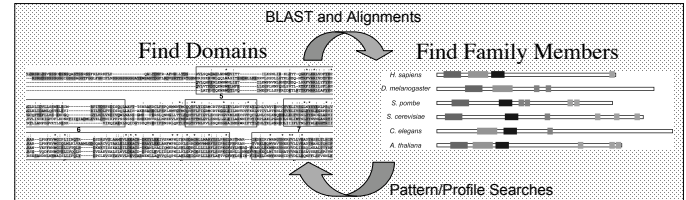


WIBR Bioinformatics Courses, © Whitehead Institute, 2004

7

Creating Protein Families

- Use domains to identify family members
 - Use a sequence to search a database and characterize a pattern/profile
 - Use a specific pattern/profile to identify homologous sequences (family members)



WIBR Bioinformatics Courses, © Whitehead Institute, 2004

8

Family Database Resources

- **Curated Databases***
 - Proteins are placed into families with which they share a specific sequence pattern
- **Clustering Databases***
 - Sequence similarity-based without the prior knowledge of specific patterns
- **Derived Databases***
 - Pool other databases into one central resource
- **Search and Browse**

*(Higgins, 2000)

WIBR Bioinformatics Courses, © Whitehead Institute, 2004

9

Curated Family Databases

- **Pfam** (<http://pfam.wustl.edu/hmmsearch.shtml/>) **
 - Uses manually constructed seed alignments and PSSM to automatically extract domains
 - db of protein families and corresponding profile-HMMs of prototypic domains
 - Searches report e-value and bits score
- **Prosite** (<http://www.expasy.ch/tools/scanprosite/>)
 - Hit or Miss -> no stats
- **PRINTS** (<http://www.bioinf.man.ac.uk/fingerPRINTScan/>)

Pfam HMM search results, glocal+local alignments merged (Pfam_ls+Pfam_fs)

[Go here for an explanation of the format of the results]

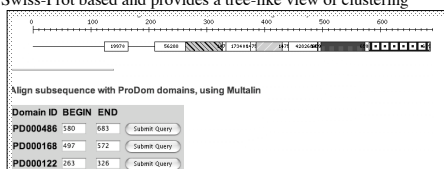
Model	Seq-from	Seq-to	HMM-from	HMM-to	Score	E-value	Alignment	Description
GTP_RFTU	258	483	1	298	315.7	5.5e-92	glocal	Elongation factor Tu GTP binding domain
GTP_RFTU_D2	502	570	1	75	46.1	8e-11	glocal	Elongation factor Tu domain 2
GTP_RFTU_D3	576	684	1	112	142.9	6.1e-40	glocal	Elongation factor Tu C-terminal domain

WIBR Bioinformatics Courses, © Whitehead Institute, 2004

10

Clustering Family Databases

- Search a database against itself and cluster similar sequences into families
- **ProDom** (<http://prodes.toulouse.inra.fr/prodom/current/html/home.php>)
 - Searchable against MSAs and consensus sequences
- **Protomap** (<http://protomap.cornell.edu/>)
 - Swiss-Prot based and provides a tree-like view of clustering

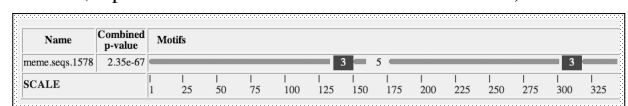


WIBR Bioinformatics Courses, © Whitehead Institute, 2004

11

Derived Family Databases

- Databases that utilize protein family groupings provided by other resources
- **Blocks** - Search and Make (<http://blocks.flirc.org/blocks/>)
 - Uses Protomap system for finding blocks that are indicative of a protein family (GIBBS/MOTIF)
- **Proclass** (<http://pir.georgetown.edu/gfserver/proclass.html>)
 - Combines families from several resources using a neural network-based system (relationships)
- **MEME** (<http://meme.sdsc.edu/meme/website/intro.html>)



WIBR Bioinformatics Courses, © Whitehead Institute, 2004

12

Syllabus

- Protein Families
 - Identifying Protein Domains
 - Family Databases & Searches
- Searching for Family Members
 - Pattern Searches
 - Patscan
 - Profile Searches
 - PSI-BLAST/HMMER2

WIBR Bioinformatics Courses, © Whitehead Institute, 2004

13

Searching Family Databases

- BLAST searches provide a great deal of information, but it is difficult to select out the important sequences (listed by score, not family)
- Family searches can give an immediate indication of a protein's classification/function
- Use Family Database search tools to identify domains and family members

WIBR Bioinformatics Courses, © Whitehead Institute, 2004

14

Patterns & Profiles

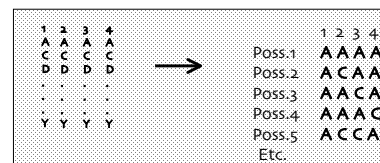
- Techniques for searching sequence databases to uncover common domains/motifs of biological significance that categorize a protein into a family
- Pattern** - a deterministic syntax that describes multiple combinations of possible residues within a protein string
- Profile** - probabilistic generalizations that assign to every segment position, a probability that each of the 20 aa will occur

WIBR Bioinformatics Courses, © Whitehead Institute, 2004

15

Discovery Algorithms

- Pattern Driven Methods
 - Enumerate all possible patterns in solution space and try matching them to a set of sequences

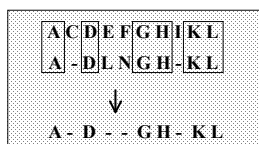


WIBR Bioinformatics Courses, © Whitehead Institute, 2004

16

Discovery Algorithms

- Sequence Driven Methods
 - Build up a pattern by pair-wise comparisons of input sequences, storing positions in common, removing positions that are different



WIBR Bioinformatics Courses, © Whitehead Institute, 2004

17

Pattern Building

- Find patterns like "pos1 xx pos2 xxxx pos3"
 - Definition of a non-contiguous motif

1.	C Y D - - C A F T L R Q S A V M H K H A R E H
2.	C A T Y - C R T A I D T V K N S L K H H S A H
3.	C W D G G C G I S V E R M D T V H K H D T V H
4.	C Y C - - C S D H M K K D A V E R M H K K D H
5.	C N M F - C M P I F R Q N S L A R E H E R M H
6.	C L N N T C T A F W R Q K K D D T V H N S L H
C xxxx C xxxx [LIVMFV] xxxxxxxx H xxxxx H	

Define/Search A Motif <http://us.expasy.org/tools/scanprosite/>

WIBR Bioinformatics Courses, © Whitehead Institute, 2004

18

Pattern Properties

- **Specification**
 - a single residue K, set of residues (KPR), exclusion {KPR}, wildcards X, varying lengths x(3,6) -> variable gap lengths
- **General Syntax**
 - C-x(2,4)-C-x(3)-[LIVMFYWC]-x(8)-H-x(3,5)-H
- **Patscan Syntax** (<http://jura.wi.mit.edu/bio/education/bioinfo-mini/seq>)
 - C 2...4 C 3...3 any(LIVMFYWC) 8...8 H 3...5 H
- **Pattern Database Searching**
 - `%scan_for_matches -p pattern_file <nr> output_file`

Sequence Pattern Concerns

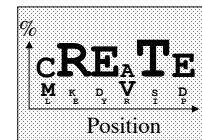
- Pattern descriptors must allow for approximate matching by defining an acceptable distance between a pattern and a potential hit
 - Weigh the sensitivity and specificity of a pattern
- What is the likelihood that a pattern would randomly occur?

Sequence Profiles

- **Consensus** - mathematical probability that a particular aa will be located at a given position
- **Probabilistic** pattern constructed from a MSA
- Opportunity to assign penalties for insertions and deletions, but not well suited for variable gap lengths
- **PSSM** - (Position Specific Scoring Matrix)
 - Represents the sequence profile in tabular form
 - Columns of weights for every aa corresponding to each column of a MSA

Profile Analysis

- Perform global MSA on group of sequences
- Move highly conserved regions to smaller MSAs
- Generate scoring table with log odds scores
 - Each column is independent
 - Average Method: profile matrix values are weighted by the proportion of each amino acid in each column of MSA
 - Evolutionary Method: calculate the evolutionary distance (Dayhoff model) required to generate the observed amino acid distribution



PSSM Example

1. I T I S
2. T D L S
3. V T M G
4. I T I G
5. V G F S
6. I E L T
7. T T T S
8. I T L S

(i.e. Distribution of aa in an MSA column)

← Target sequences

Resulting Consensus: I T L S

PSSM

P O S																				
	A	C	D	E	F	G	H	I	K	L	M	N	P	Q	R	S	T	V	W	Y
1	8	-2	5	4	5	5	-4	24	0	15	13	1	1	1	-7	2	22	21	-18	-6
2	13	-5	24	18	-18	19	7	1	7	-7	-4	14	11	10	-1	9	29	3	-28	-14
3	5	-5	3	4	13	4	2	8	-4	14	12	8	-5	0	-10	0	10	10	-1	5
4	17	17	13	10	-12	29	-5	-5	6	-14	-9	12	10	0	-2	34	19	1	-8	-15

PSSM Properties

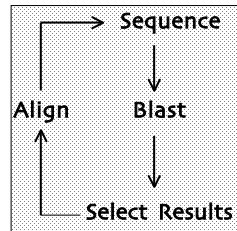
- Score-based sequence representations for searching databases
- Goal
 - Limit the diversity in each column to improve reliability
- Problems
 - Differing length gaps between conserved positions (unlike patterns)

PSI-BLAST Implementation

- **PSI-BLAST**

<http://www.ncbi.nlm.nih.gov/BLAST/>

- Start with a sequence, BLAST it, align select results to query sequence, estimate a profile with the MSA, search DB with the profile - constructs PSSM
- Iterate until process stabilizes
- Focus on domains, not entire sequences
- Greatly improves sensitivity



PSI-BLAST Sample Output

Sequences with E-value WORSE than threshold

gi19629055 ref NP_044074.1	(NC_001721) NC123R [Mollusca contig...	37	0.16
gi16176554 db AA035499.2	(S79774) bile salt-dependent lipase: B...	36	0.25
gi1450277 ref NP_001798.1	(NM_001807) carboxyl ester lipase (D...	35	0.86
gi1231628 ref NP_003518	hLIP_HUMAN Bile-salt-activated lipase precu...	35	0.89
gi15242929 ref NP_208612.1	(NM_125189) putative protein [Arabid...	34	1.1
gi19769529 db AB010925.1	(AB024029) gene_id:K21L19.3-unknown p...	34	1.3
gi1804807 db AA052014.1	(M85201) cholesterol esterase [Homo sap...	33	1.8
gi110706 ref NP_173109	hLIP_HUMAN Chromosomal replication initiat...	32	4.6
gi1126679 ref NP_110110	hLIP_HUMAN GALECTIN-3 (GALACTOSE-SPECIFIC LE...	32	4.9
gi152851 ref NP_001000	(X16074) L-34 protein (AA 1-264) [Mus sp.]	32	5.0
gi153907 ref NP_110110	hLIP_HUMAN Lactose-binding lectin Mac-2 - mouse	32	5.0
gi1387111 ref NP_001000	(J03723) carbohydrate binding protein 3...	32	5.4
gi1950427 ref NP_062019.1	(NM_019146) bassoon [Pattus norvegic...	32	5.5

HMM Building

- **Hidden Markov Models** are Statistical methods that consider all the possible combinations of matches, mismatches, and gaps to generate a consensus (Higgins, 2000)
- Sequence ordering and alignments are not necessary at the onset (but in many cases alignments are recommended)
- Ideally use at least 20 sequences in the training set to build a model
- Calibration prevents over-fitting training set (i.e. Ala scan)
- Generate a model (profile/PSSM), then search a database with it

HMM Implementation

- **HMMER2** (<http://hmmerr.wustl.edu/>)
 - Determine which sequences to include/exclude
 - Perform alignment, select domain, excise ends, manually refine MSA (pre-aligned sequences better)
 - Build profile
 - `%hmmbuild [-options] <hmmfile> <alignment file>`
 - Calibrate profile (re-calc. Parameters by making a random db)
 - `%hmmcalibrate [-options] <hmmfile>`
 - Search database
 - `%hmmsearch [-options] <hmmfile> <database file> > out`

HMMER2 Output

- Hmmssearch returns e-values and bits scores
- Repeat process with selected results
 - Unfortunately need to extract sequences from the results and manually perform MSA before beginning next round of iteration

```
HMMER 2.2g (August 2001)
Copyright (C) 1992-2001 HHMI/Washington University School of Medicine
Freely distributed under the GNU General Public License (GPL)

HMM file: pfam_hmm [Hydrolase]
Sequence database: /cluster/00/Data/r
per-sequence score cutoff: [none]
per-domain score cutoff: [none]
per-sequence E-value cutoff: <= 10
per-domain E-value cutoff: [none]

Query HMM: Hydrolase
Accession: PF00702
Description: halocidal dehalogenase-like hydrolase
[HMM has been calibrated: E-values are empirical estimates]

Scores for complete sequences (score includes all domains):
Sequence      Description      Score      E-value  N
-----
gi1613126|ref|NP_417844.1|  phosphoglycolat 168.4  2.9e-45  1
gi24114648|ref|NP_209158.1|  phosphoglycolat 167.8  4.2e-45  1
gi15803888|ref|NP_289924.1|  phosphoglycolat 167.8  4.2e-45  1
gi26249979|ref|NP_756019.1|  Phosphoglycolat 166.4  1.1e-44  1
```

Patterns vs. Profiles

- **Patterns**
 - Easy to understand (human-readable)
 - Account for different length gaps
- **Profiles**
 - Sensitivity, better signal to noise ratio
 - Teachable

Domain ID & Searching

- **Family/Domain Search**
 - <http://pfam.wustl.edu>
- **Pattern Search**
 - `scan_for_matches (Patscan)`
 - `scan_for_matches -p pattern_file <[cluster/db0/Data/yeast.aa >output_file`
any(CF) any(HF) any(K) 1..1 any(L) 4..4 any(AS) 3..3 any(L) 3..3 any(GA) 3..3 G 1..1 any(YF) 1..1 any(L) R
- **Profile Search**
 - **HMMER2**
 - `hmmdb [-options] <hmmfile output> <alignment file>`

[illegible]

Exercises

- Use PFAM to identify domains within your sequence
- Scan your sequences with ProSite to find a pattern to represent the domain
- Use the ProSite pattern to search the non-redundant db
- Use PSI-BLAST to build a sequence profile and search the non-redundant db

References

- Bioinformatics: Sequence and genome Analysis. David W. Mount. CSHL Press, 2001.
- Bioinformatics: A Practical Guide to the Analysis of Genes and Proteins. Andreas D. Baxevanis and B.F. Francis Ouellete. Wiley Interscience, 2001.
- Bioinformatics: Sequence, structure, and databanks. Des Higgins and Willie Taylor. Oxford University Press, 2000.