

8th RepeatExplorer Workshop on the Application of Next Generation Sequencing to Repetitive DNA Analysis

Welcome !



*Institute of Plant Molecular Biology,
Biology Centre CAS
České Budějovice, Czech Rep.*



Acknowledgments



Biology Centre CAS, České Budějovice
Institute of Plant Molecular Biology

Laboratory of Molecular
Cytogenetics

Petr Novák
Pavel Neumann
Iva Fuková
Laura A. Robledillo
Nina Hošťáková
Ludmila Oliveira

Renata Marecová
Andrea Koblížková
Vlasta Tetourová
Jana Látalová
Tihana Vondrak



Programme

Tuesday (May 21)

9:15 - **Principles and applications of graph-based repeat clustering** (J. Macas)

10:00 - **RepeatExplorer pipeline** (P. Novák)

10:30 - 11:00 *Coffee break*

11:00 - **REXdb database and transposon classification using conserved protein domains** (P. Neumann)

11:30 - **Using RepeatExplorer output for repeat annotation and quantification** (J. Macas)

11:50 - **Additional RE tools** (P. Novák)

12:30 - 13:30 *Lunch*

13:30 - (18:00) **Practical training I**

- design of sequencing and repeat analysis experiments
- introduction to Galaxy environment
- quality control and pre-processing of NGS reads, dealing with various read formats
- setting up clustering analysis
- comparative clustering of multiple samples

19:00 → : Dinner at “CITYgastro”

Wednesday (May 22)

8:30 -12:30: **Short talks**

- Steven Dodsworth
- Lucia Campos-Dominguez
- Santelmo Vasconcelos
- Sidonie Bellot

Coffee break

- Tony Heitkam
- Vildana Suljevic & Hanna Schneeweiss
- Ludwig Mann
- Aleš Kovařík

12:30 - 13:30 *Lunch*

13:30 - (18:00) **Practical training II**

- identification of satellite DNA using TAREAN
- understanding RepeatExplorer output
- cluster annotation and repeat composition of the genome
- comparative clustering of multiple species – data interpretation
- repeat quantification (principles, sensitivity and reproducibility)
- design of hybridization probes based on RE

Thursday (May 23)

8:30 -12:30: **Short talks**

- Aretuza Sousa
- Kathrin Seibt
- Vratislav Peška
- Jasna Puizina

Coffee break

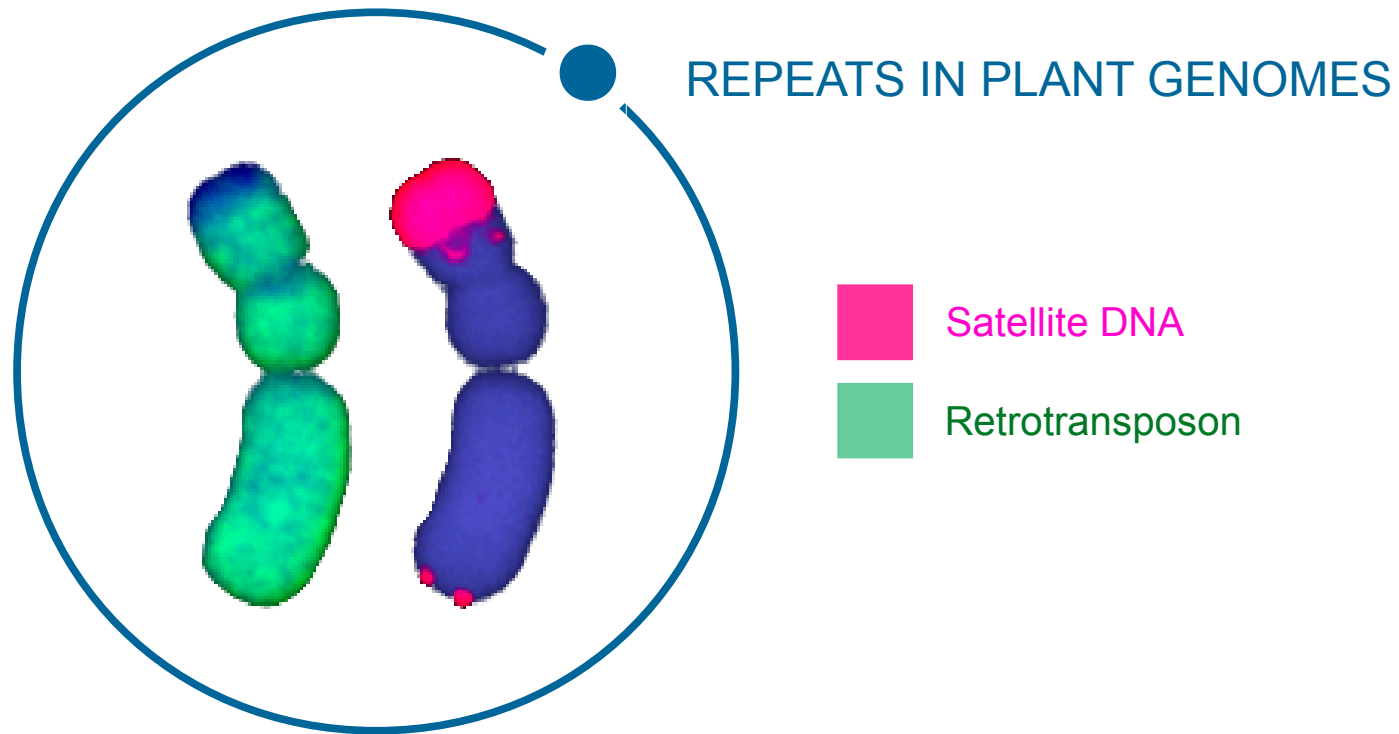
- Mariana Baez
- Nicola Schmidt
- Mark Johnston
- Horacio Naveira

12:30 - 13:30 *Lunch*

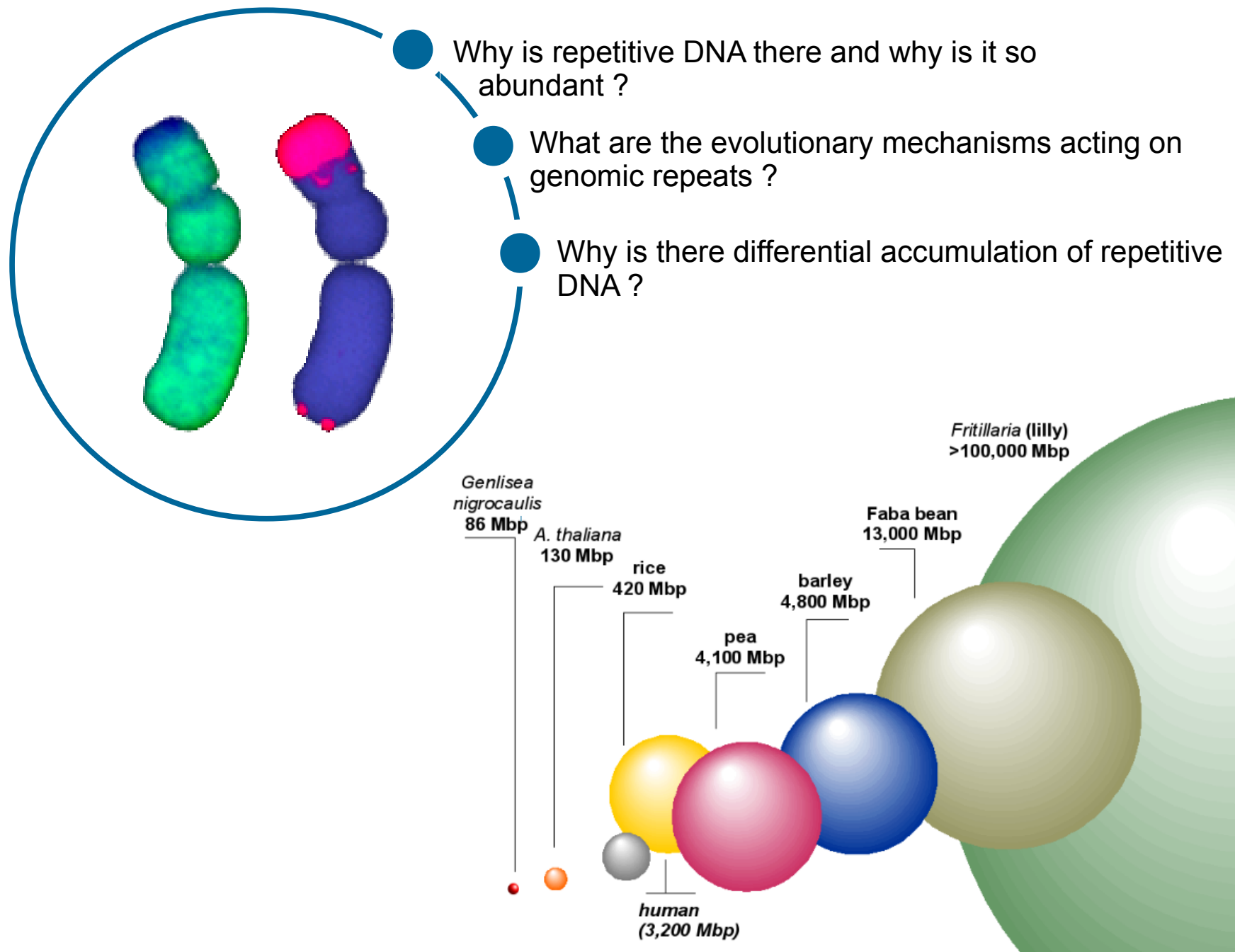
13:30 – (18:00) - **Practical training III**

- combining repeat clustering with ChIP-seq data
- identification and phylogenetic analysis of retrotransposon protein domains
- SeqGrapheR – visualization and annotation of the cluster graphs
- *advanced topics, troubleshooting*

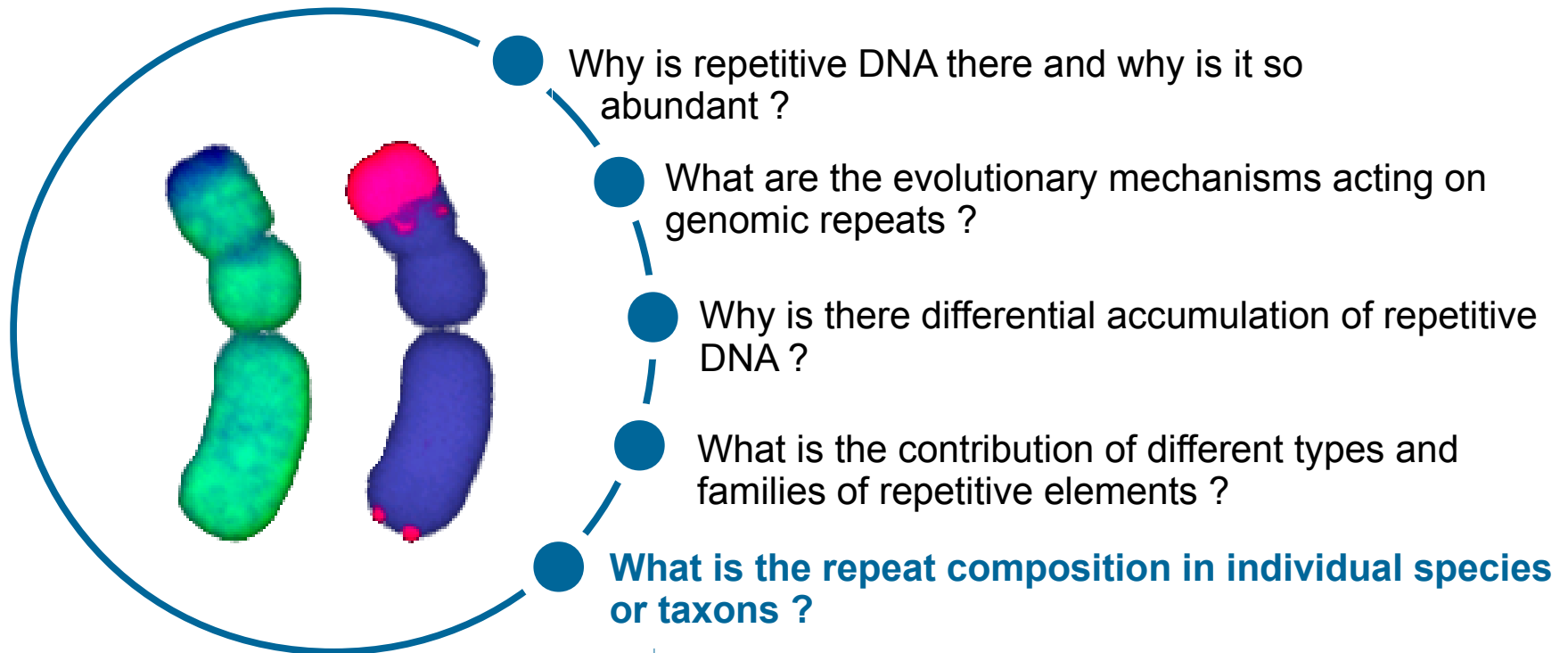
Principles and applications of graph-based repeat clustering



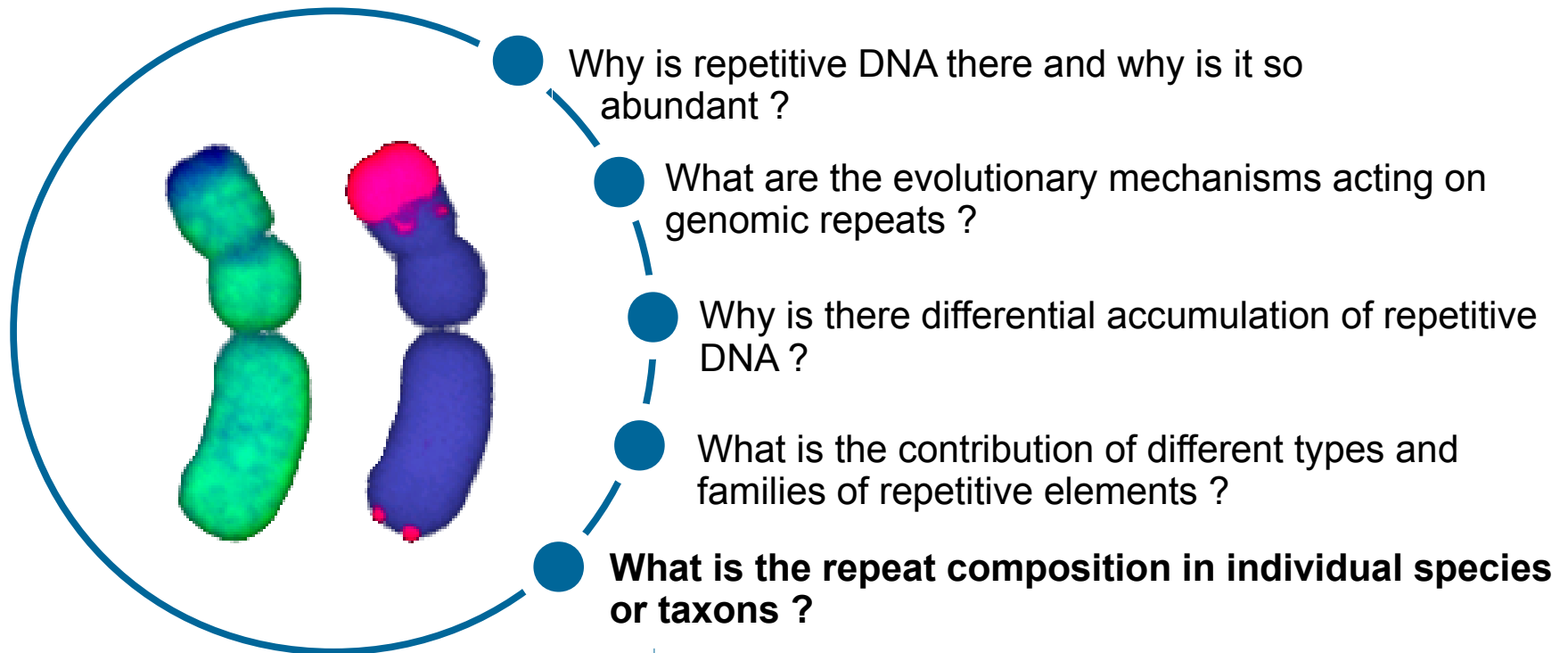
Repeats in plant genomes



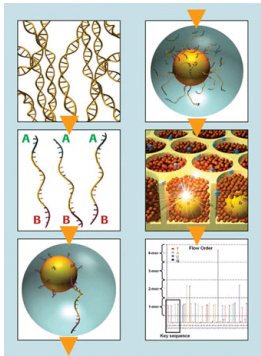
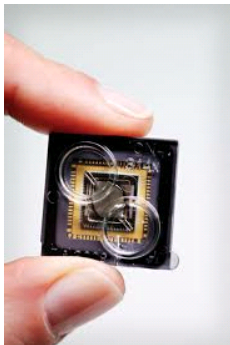
Repeats in plant genomes



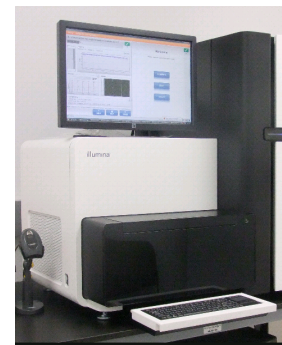
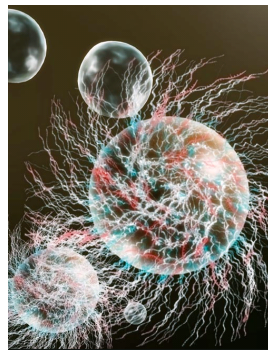
The challenge of repeat identification in complex genomes



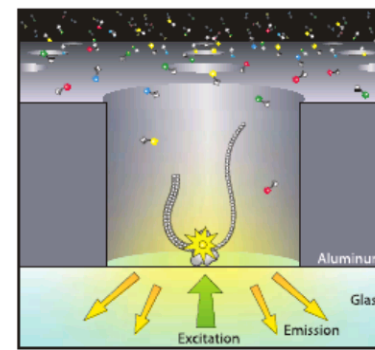
Next generation sequencing: getting sequence data is no longer a limiting factor



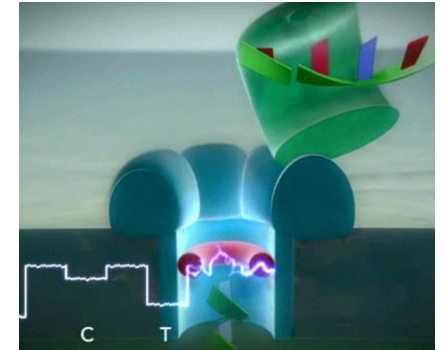
454/Roche



Illumina



Pacific Biosciences



Oxford Nanopore

Principles and tools for global repeat identification

		Assembled genome or long contigs (BACs etc.)	
		Y E S	N O
Reference database	Y E S		
	N O		

Repeat identification in sequence data

- Pure algorithmic (evaluates “repetitiveness” of substrings in DNA sequence)
- Signature or feature-based (searches for specific structures of biological importance)
- Similarity to known repeats

Principles and tools for global repeat identification

		Assembled genome or long contigs (BACs etc.)	
		Y E S	N O
Reference database	Y E S	RepeatMasker / RepBase • <i>first “model” species</i> • <i>extensive resources (+ wet lab)</i>	• <i>model-related taxa</i> • <i>often shotgun, “assembled”</i>
	N O		

Repeat identification in sequence data

- Pure algorithmic (evaluates “repetitiveness” of substrings in DNA sequence)
- Signature or feature-based (searches for specific structures of biological importance)
- Similarity to known repeats (uses curated databases of known repeats)

Principles and tools for global repeat identification

RepeatMasker

(A.F.A. Smit, R. Hubley & P. Green)

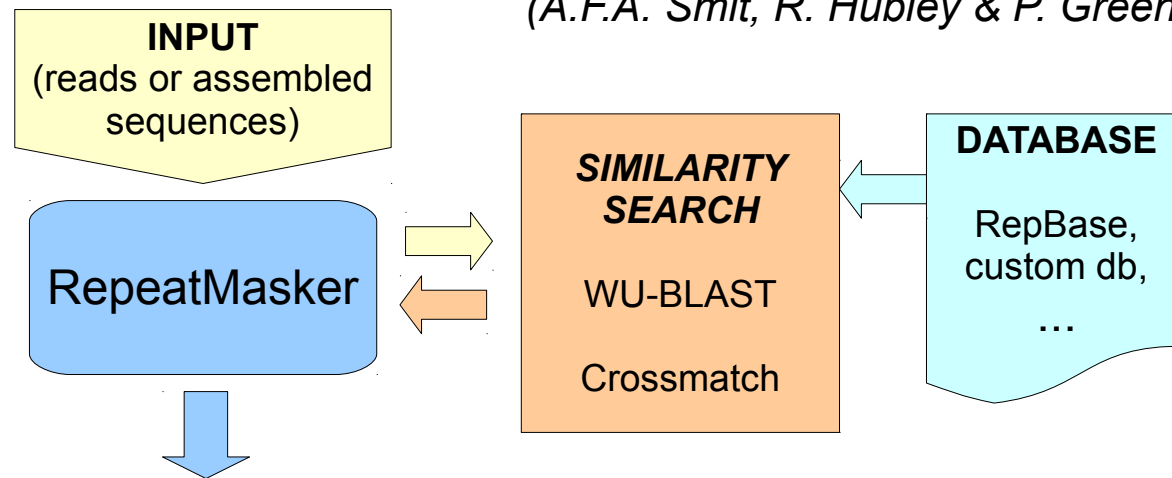


table of hits

SW score	perc div.	perc del.	perc ins.	query sequence	position begin	position end	in query (left)	matching repeat	repeat class/family
459	20.0	1.1	0.0	LAL_1000603f	2	96	(4) +	Ogre_PS1_from_AY299398	LTR/Gypsy/Ogre
361	27.8	0.0	0.0	LAL_1000670r	1	90	(10) C	Ogre_PA_z_c700	LTR/Gypsy/Ogre_PA
449	23.5	0.0	0.0	LAL_1001438f	1	98	(2) +	PSC454_CL1Contig253	LTR/Gypsy/Ogre
491	23.5	0.0	0.0	LAL_1001438r	3	100	(0) C	PSC454_CL1Contig253	LTR/Gypsy/Ogre
487	24.0	0.0	0.0	LAL_1001478f	1	100	(0) C	PSC454_CL1Contig253	LTR/Gypsy/Ogre
398	28.0	0.0	0.0	LAL_1001478r	1	100	(0) C	PSC454_CL1Contig2773	LTR/Gypsy/Ogre
417	20.6	2.0	2.0	LAL_1002130r	2	100	(0) C	PSC454_CL1Contig253	LTR/Gypsy/Ogre
373	29.0	0.0	0.0	LAL_1002357f	1	100	(0) +	Ogre_PS1_from_AY299398	LTR/Gypsy/Ogre

```
>seq_1
NNNNNNNNATGAATGCCTGTCCCCGTCAGGGTTGTTGAGTTTTTCCCTAGC
AAATCTCCCGGGCAGAGGTTTGTGTGGATGGTTTTTCCCAGCAGGNNNNN
NNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNN
NNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNN
NNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNN
NNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNNCCGATCGTCCGTTGACGTCGGGT
TGATCAGATTTTCCCCCAGAGTCACCTCTCCGTTGGTGTGCGATAAACGAT
GAGTTTTTTCCTATGCGTCCGTTGACGTATAGCTGCATGTTCCCCAAAGAT
CATTTGTTAATGC
```

< masked
sequence

summary >
table

	number of elements*	length occupied	percentage of sequence
Retroelements	1	12014 bp	96.40 %
SINEs:	0	0 bp	0.00 %
Penelope	0	0 bp	0.00 %
LINEs:	0	0 bp	0.00 %
CRE/SLACS	0	0 bp	0.00 %
L2/CR1/Rex	0	0 bp	0.00 %
R1/LOA/Jockey	0	0 bp	0.00 %
R2/R4/NeSL	0	0 bp	0.00 %
RTE/Bov-B	0	0 bp	0.00 %
L1/CIN4	0	0 bp	0.00 %
LTR elements:	1	12014 bp	96.40 %
BEL/Pao	0	0 bp	0.00 %
Ty1/Copia	0	0 bp	0.00 %
Gypsy/DIRS1	1	12014 bp	96.40 %
Retroviral	0	0 bp	0.00 %
DNA transposons	0	0 bp	0.00 %
hobo-Activator	0	0 bp	0.00 %
Tc1-IS630-Pogo	0	0 bp	0.00 %
En-Spm	0	0 bp	0.00 %
MuDR-IS905	0	0 bp	0.00 %
PiggyBac	0	0 bp	0.00 %
Tourist/Harbinger	0	0 bp	0.00 %
Other (Mirage, P-element, Transib)	0	0 bp	0.00 %

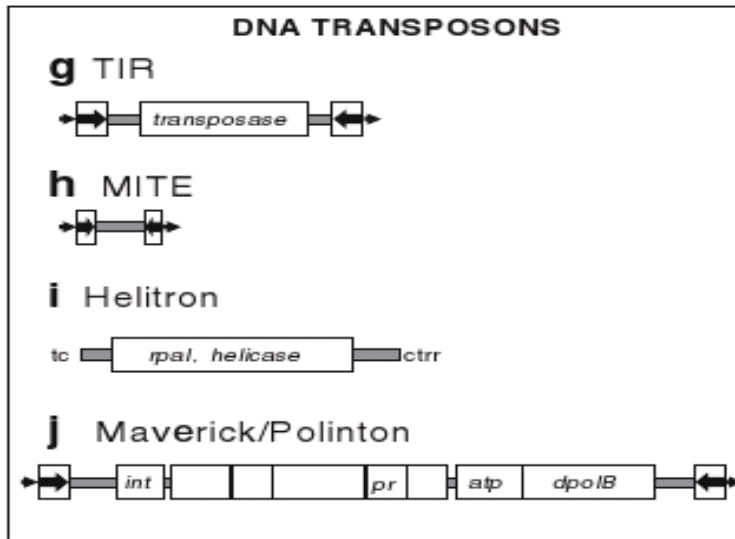
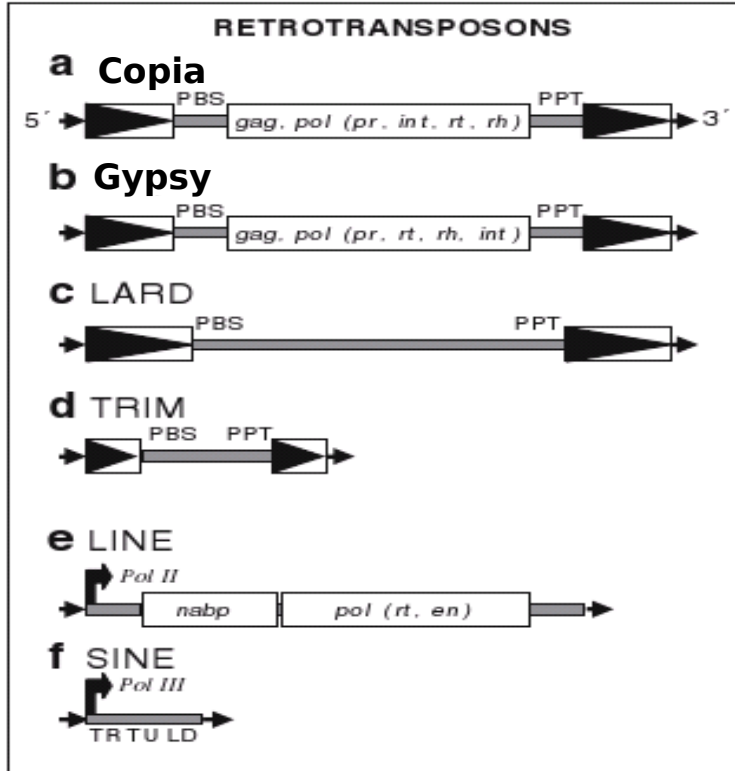
Principles and tools for global repeat identification

		Assembled genome or long contigs (BACs etc.)	
		Y E S	N O
Reference database	Y E S	RepeatMasker / RepBase • <i>first “model” species</i> • <i>extensive resources (+ wet lab)</i>	• <i>model-related taxa</i> • <i>often shotgun, “assembled”</i>
	N O	• LTR_STRUC, MITE-Hunter , etc.	

Repeat identification in sequence data

- Pure algorithmic (evaluates “repetitiveness” of substrings in DNA sequence)
- **Signature or feature-based (searches for specific structures of biological importance)**
- Similarity to known repeats (uses curated databases of known repeats)

Principles and tools for global repeat identification



Signature / structure-based approaches

LTR_STRUC (McCarthy and McDonald, 2003)

LTR_FINDER (Xu and Wang, 2007)

SINE-Finder (Wenke et al. 2011)

MITE-Hunter (Han and Wessler, 2010)

MITE Analysis Toolkit (Yang and Hall, 2003)

HelitronFinder (Du et al. 2008) [HelA helitrons in maize]

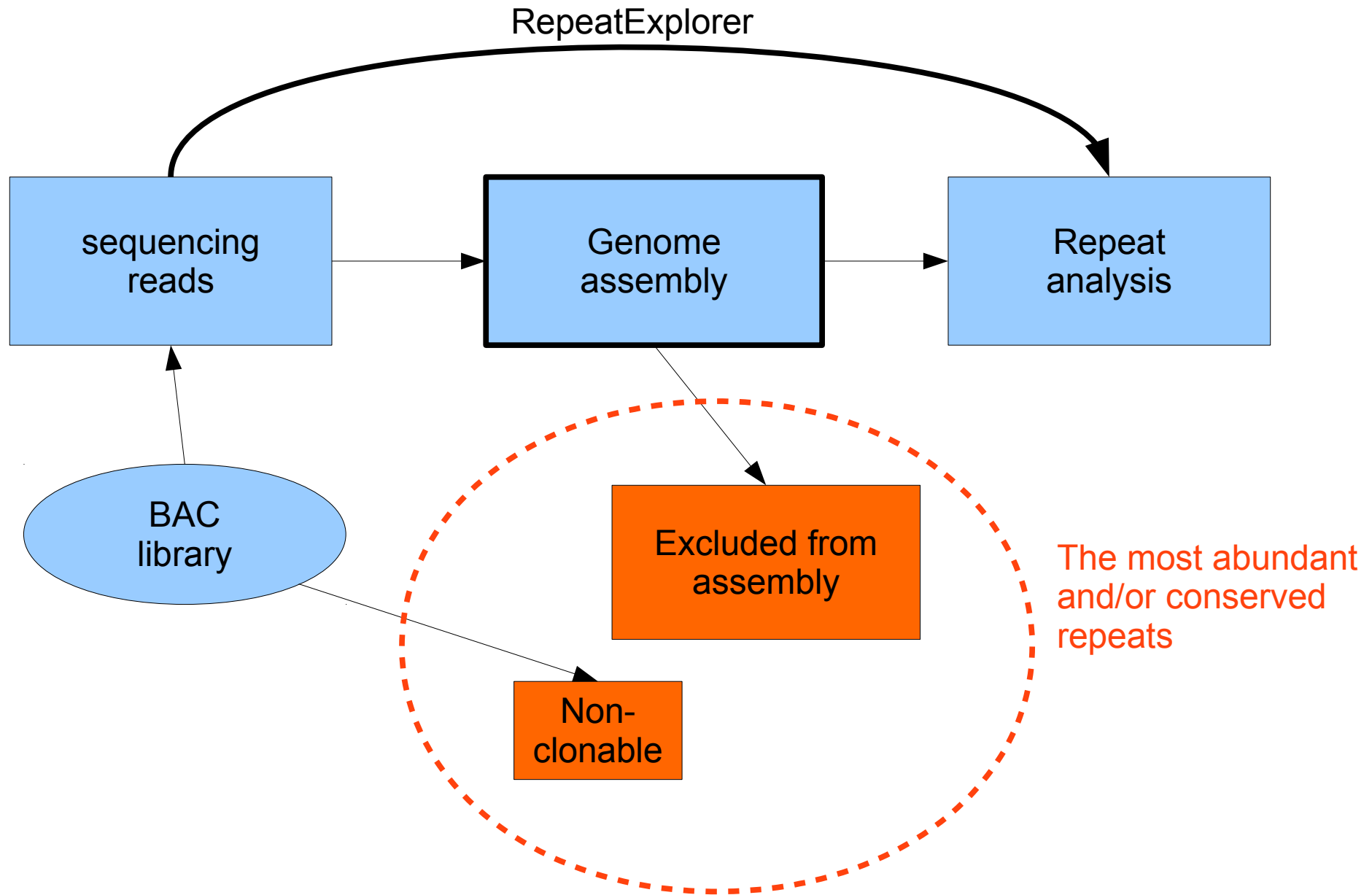
Principles and tools for global repeat identification

		Assembled genome or long contigs (BACs etc.)	
		YES	NO
Reference database	YES	RepeatMasker / RepBase • <i>first “model” species</i> • <i>extensive resources (+ wet lab)</i>	• <i>model-related taxa</i> • <i>often shotgun, “assembled”</i>
	NO	• LTR_STRUC, MITE-Hunter , etc. • REPET	<i>(using NGS reads)</i> • direct assembly (phrap,...) • clustering-based

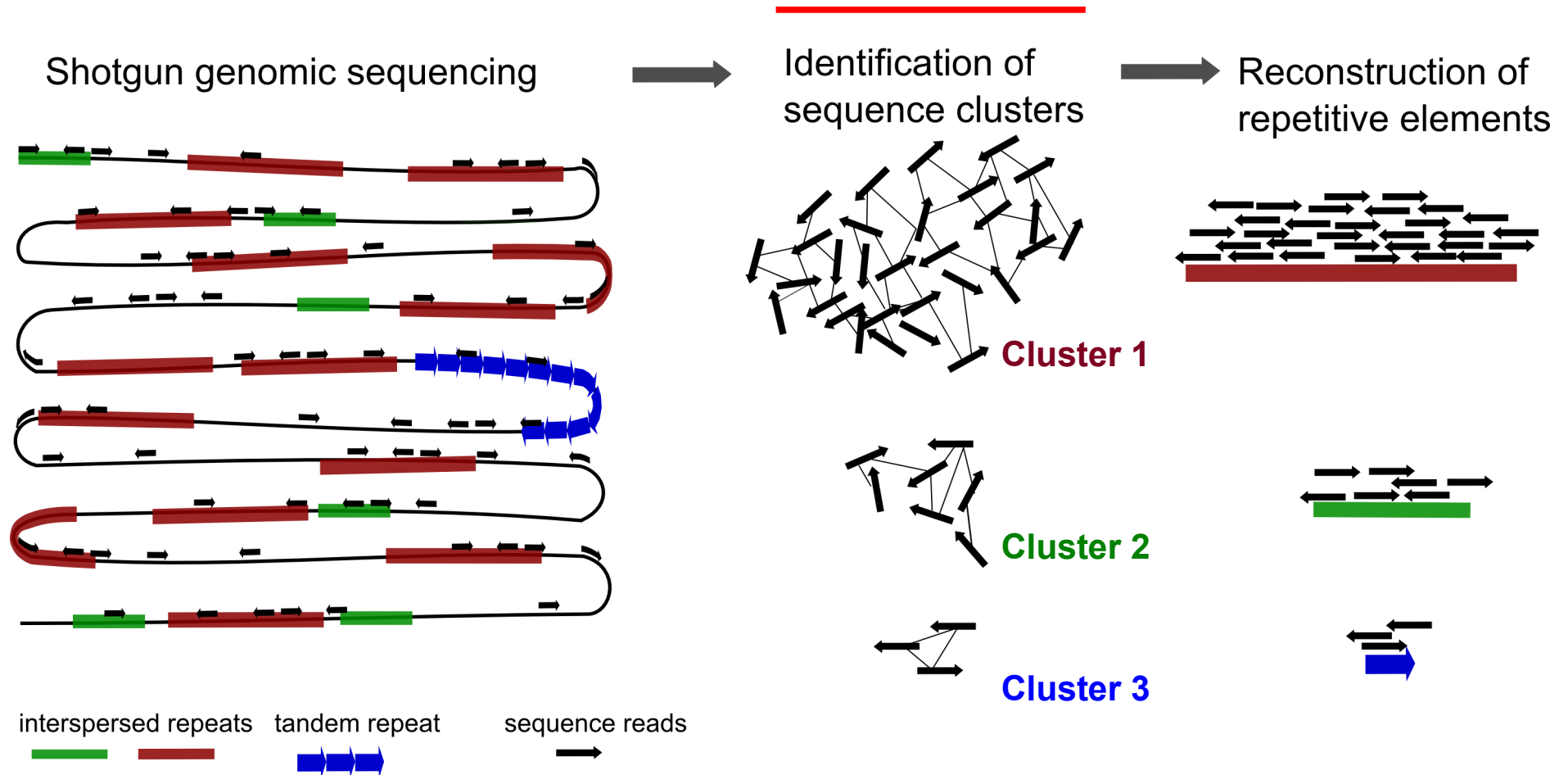
Repeat identification in sequence data

- **Pure algorithmic (evaluates “repetitiveness” of substrings in DNA sequence)**
- Signature or feature-based (searches for specific structures of biological importance)
- Similarity to known repeats (uses curated databases of known repeats)

Principles and tools for global repeat identification

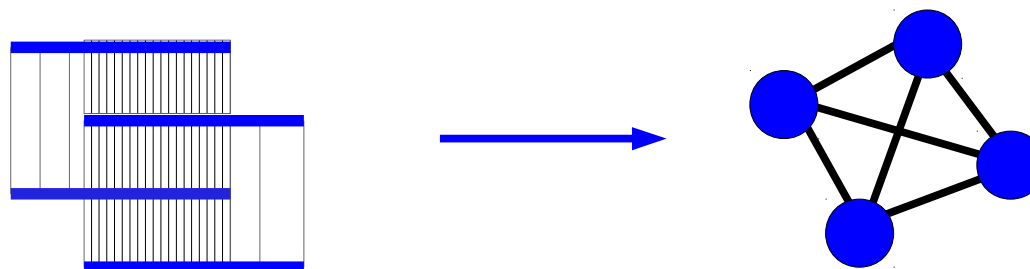
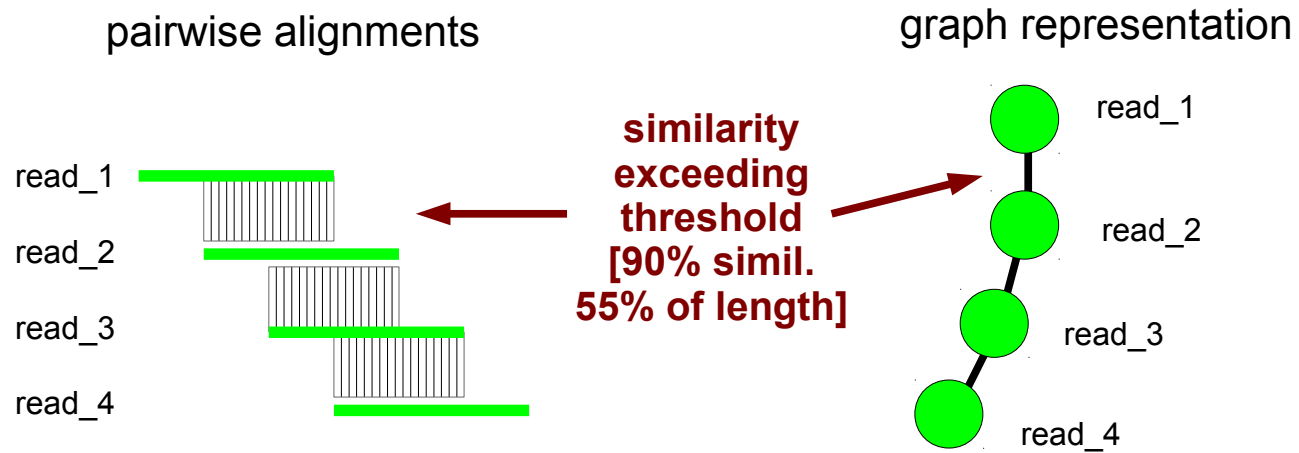


Clustering-based repeat identification

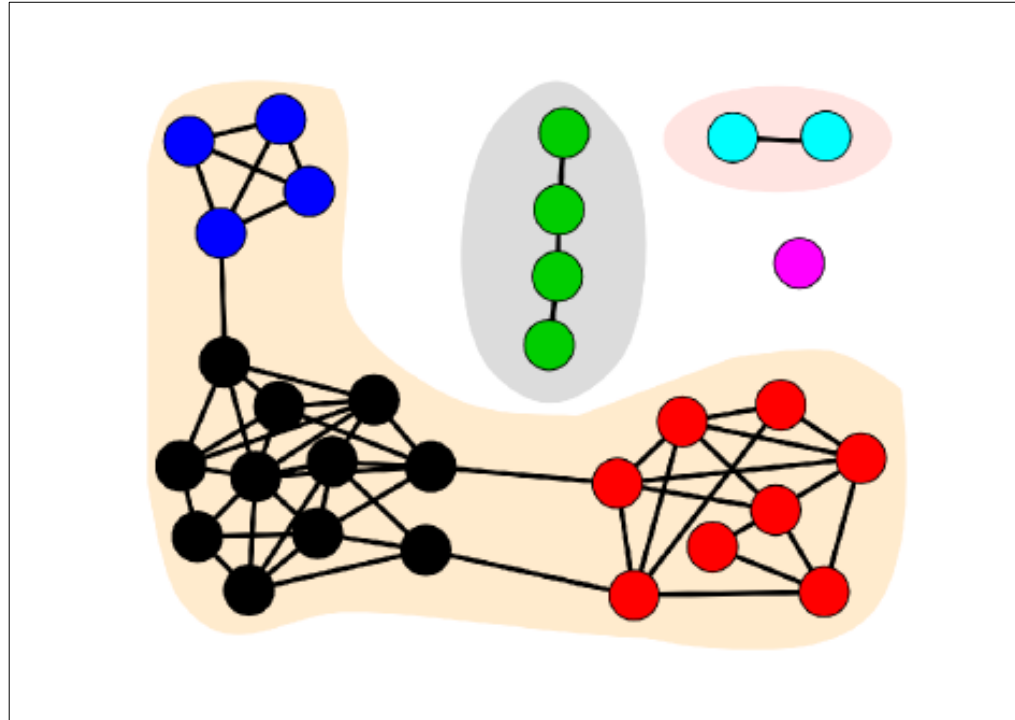


CLUSTER = a set of frequently overlapping reads = REPEAT FAMILY

Clustering-based repeat identification



Clustering-based repeat identification

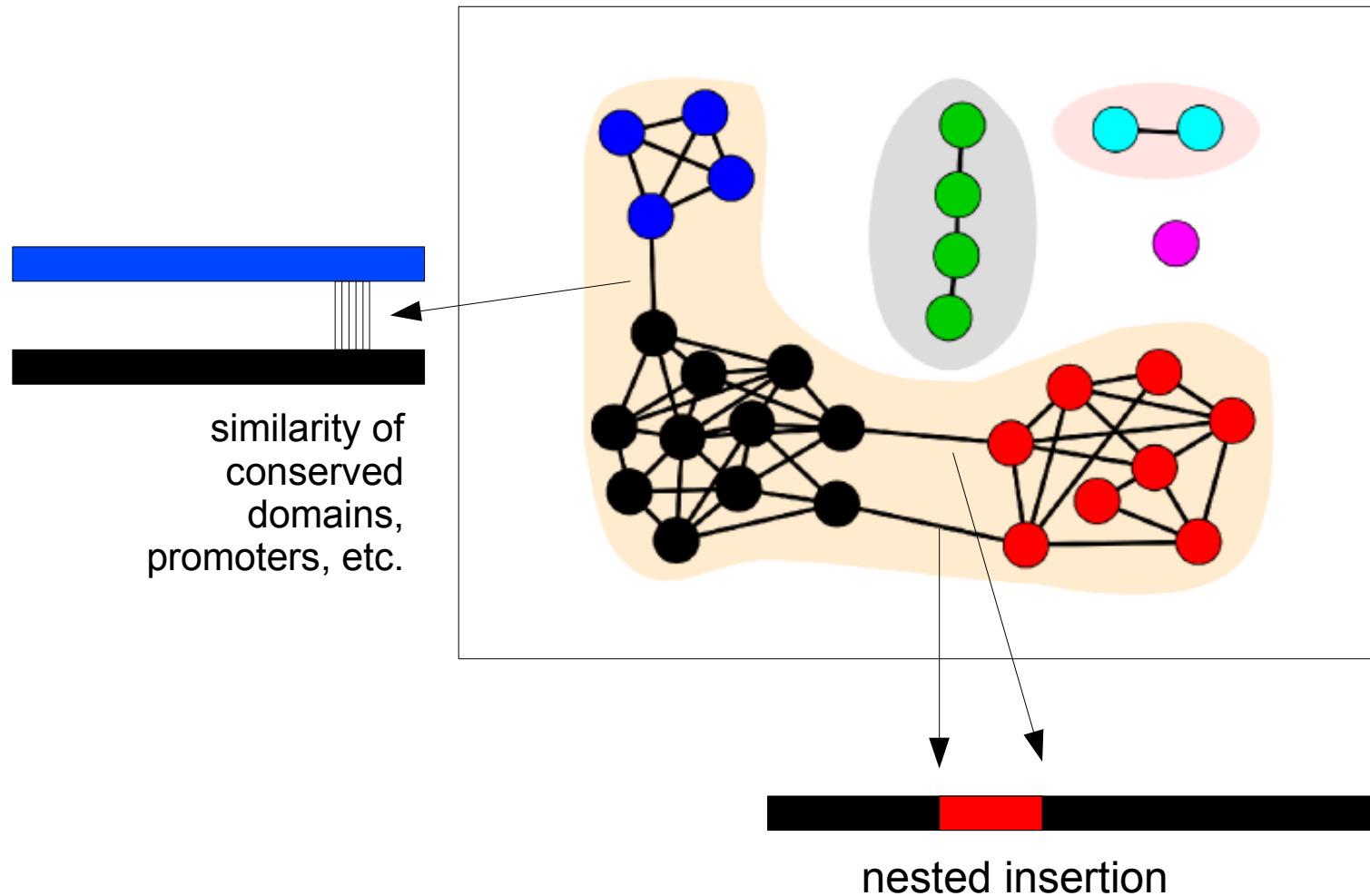


Single linkage clustering => connected components

TGICL
(TIGR Gene Indices clustering tool)
Pertea et al., 2003

Macas et al. (2007) - Repetitive DNA in the pea (*Pisum sativum* L.) genome: comprehensive characterization using 454 sequencing and comparison to soybean and *Medicago truncatula*.

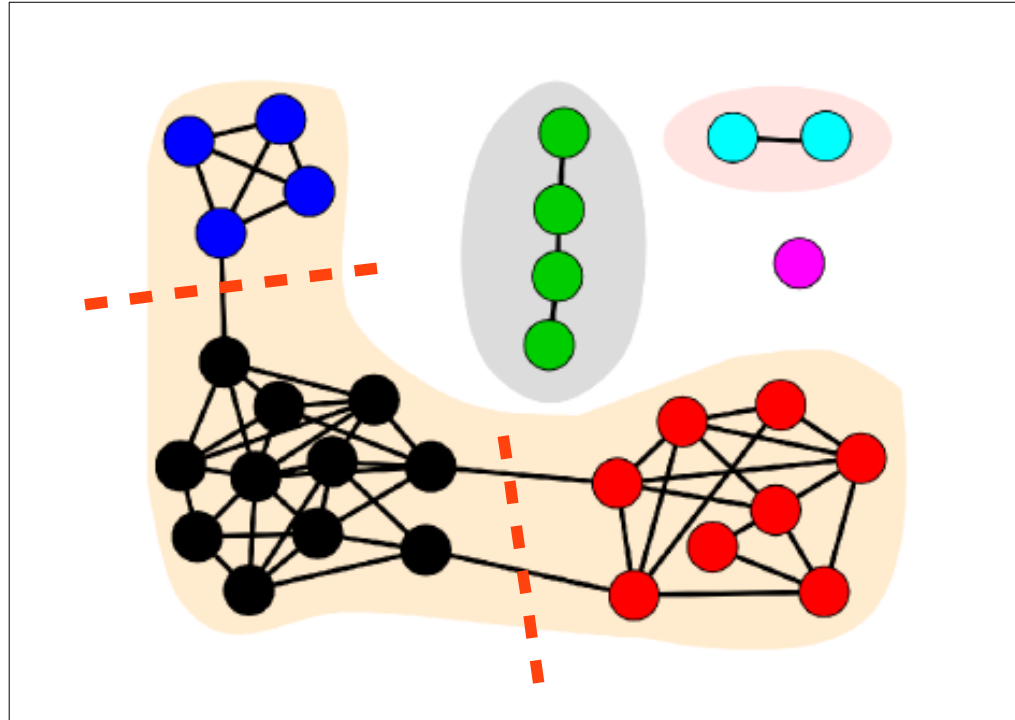
Clustering-based repeat identification



Single linkage clustering => connected components

TGICL
(TIGR Gene Indices clustering tool)
Pertea et al., 2003

Graph-based clustering

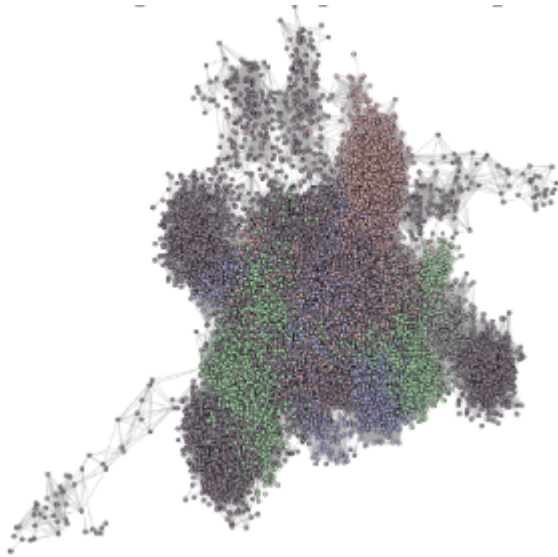


Graph-based clustering

- Sequence overlaps between the reads are transformed to a graph where the **reads** are represented as **nodes** and their **similarities** as **edges** connecting the nodes
- Graph structure is examined to detect *communities of frequently connected nodes* which are split to separate clusters



Graph-based characterization of repeat clusters

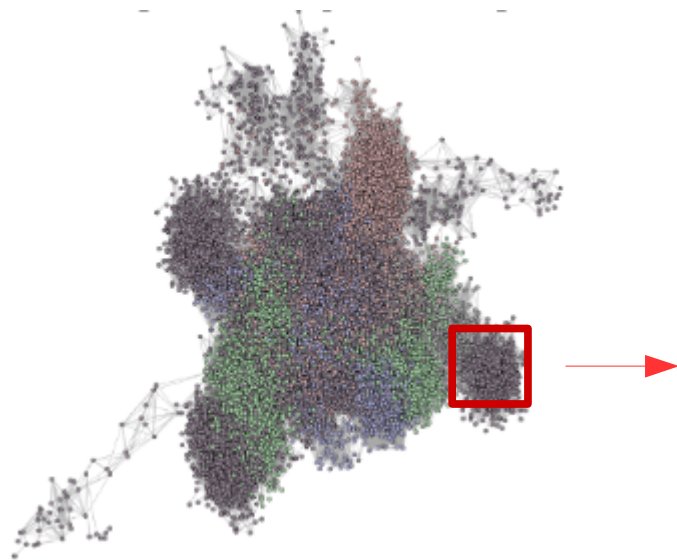


Virtual graphs used
to analyze real data
contain **up to millions**
of nodes (reads)

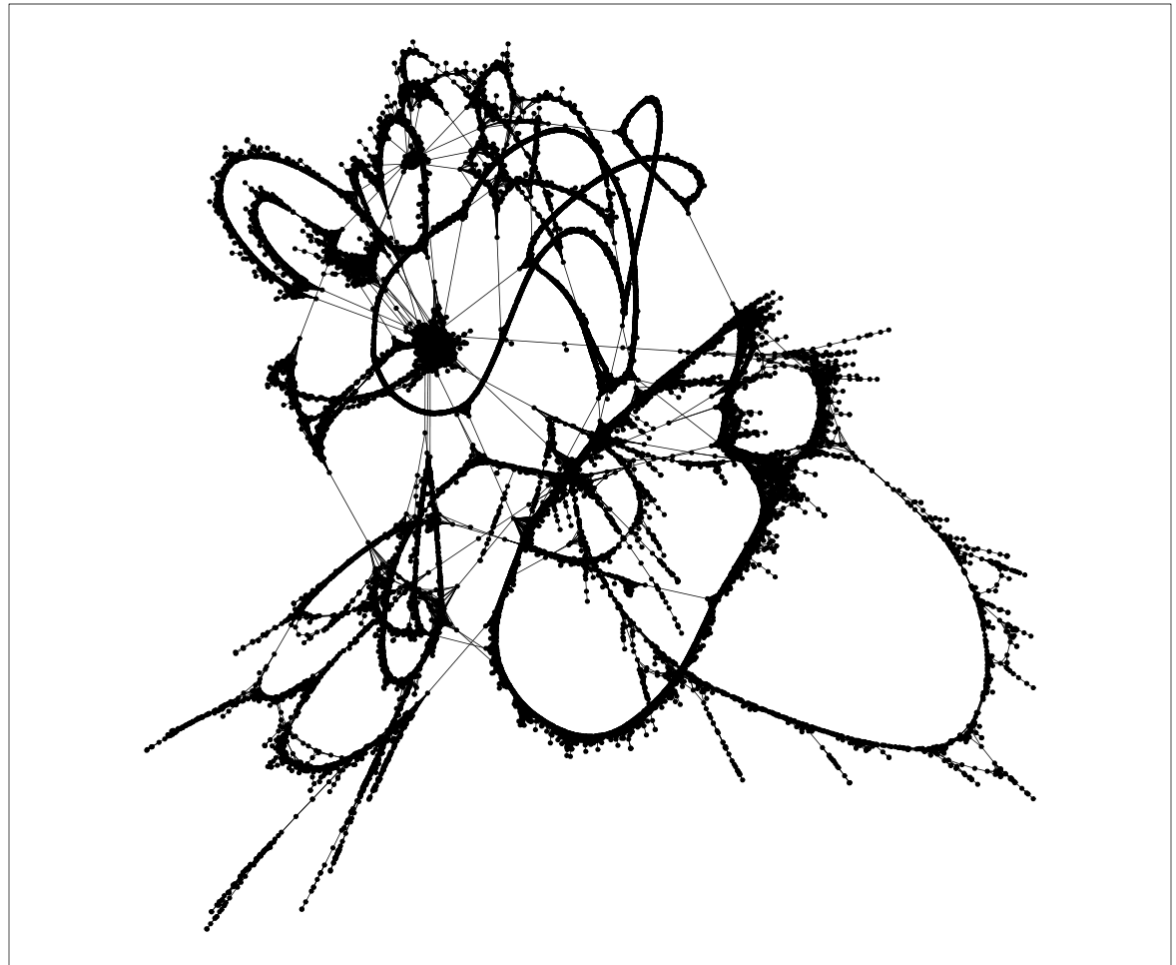
Graph-based clustering

- Sequence overlaps between the reads are transformed to a graph where the **reads** are represented as **nodes** and their **similarities** as **edges** connecting the nodes
- Graph structure is examined to detect *communities of frequently connected nodes* which are split to separate clusters

Graph-based characterization of repeat clusters



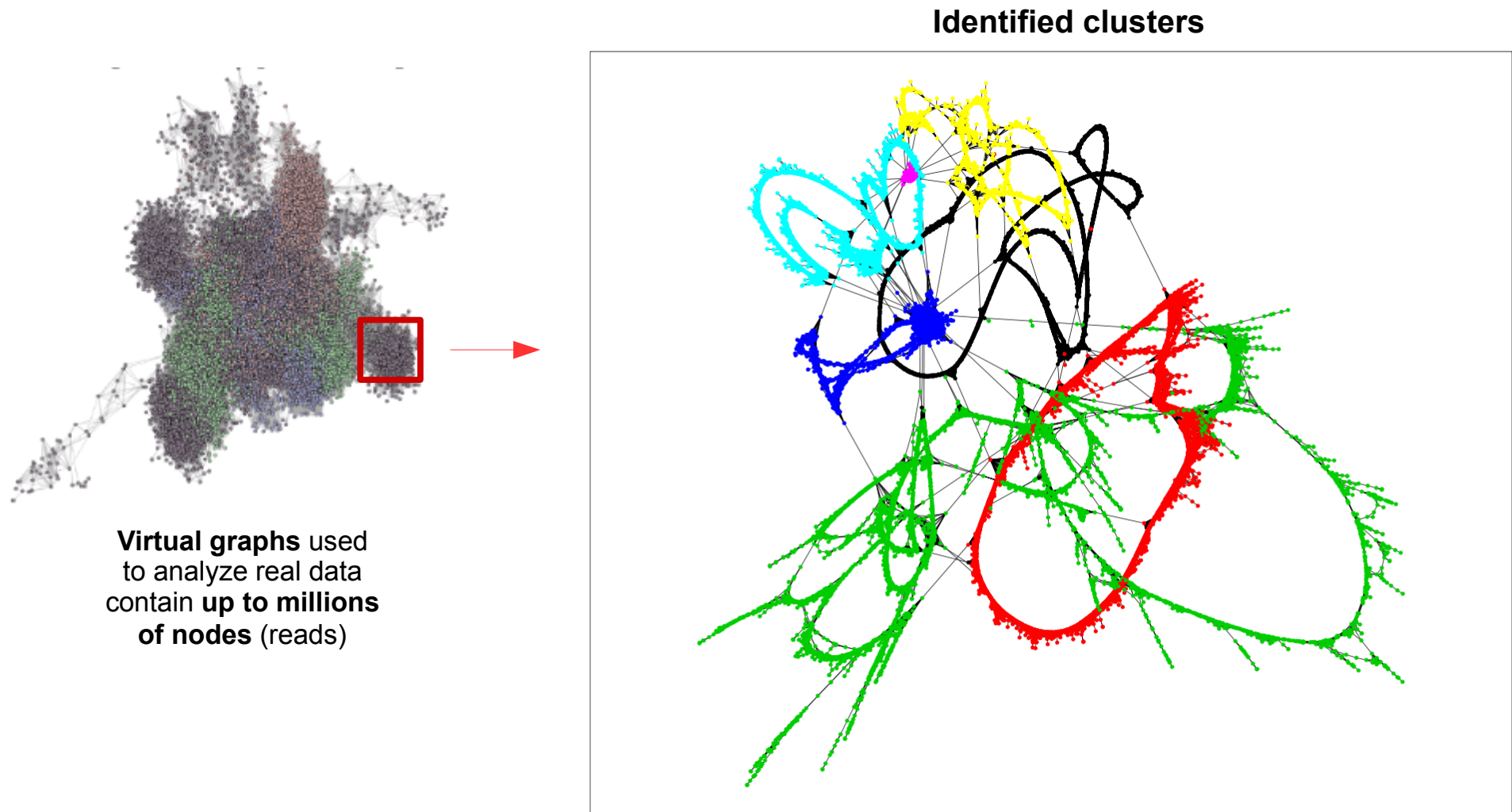
Virtual graphs used to analyze real data contain **up to millions of nodes** (reads)



Graph-based clustering

- Sequence overlaps between the reads are transformed to a graph where the **reads** are represented as **nodes** and their **similarities** as **edges** connecting the nodes
- Graph structure is examined to detect *communities of frequently connected nodes* which are split to separate clusters

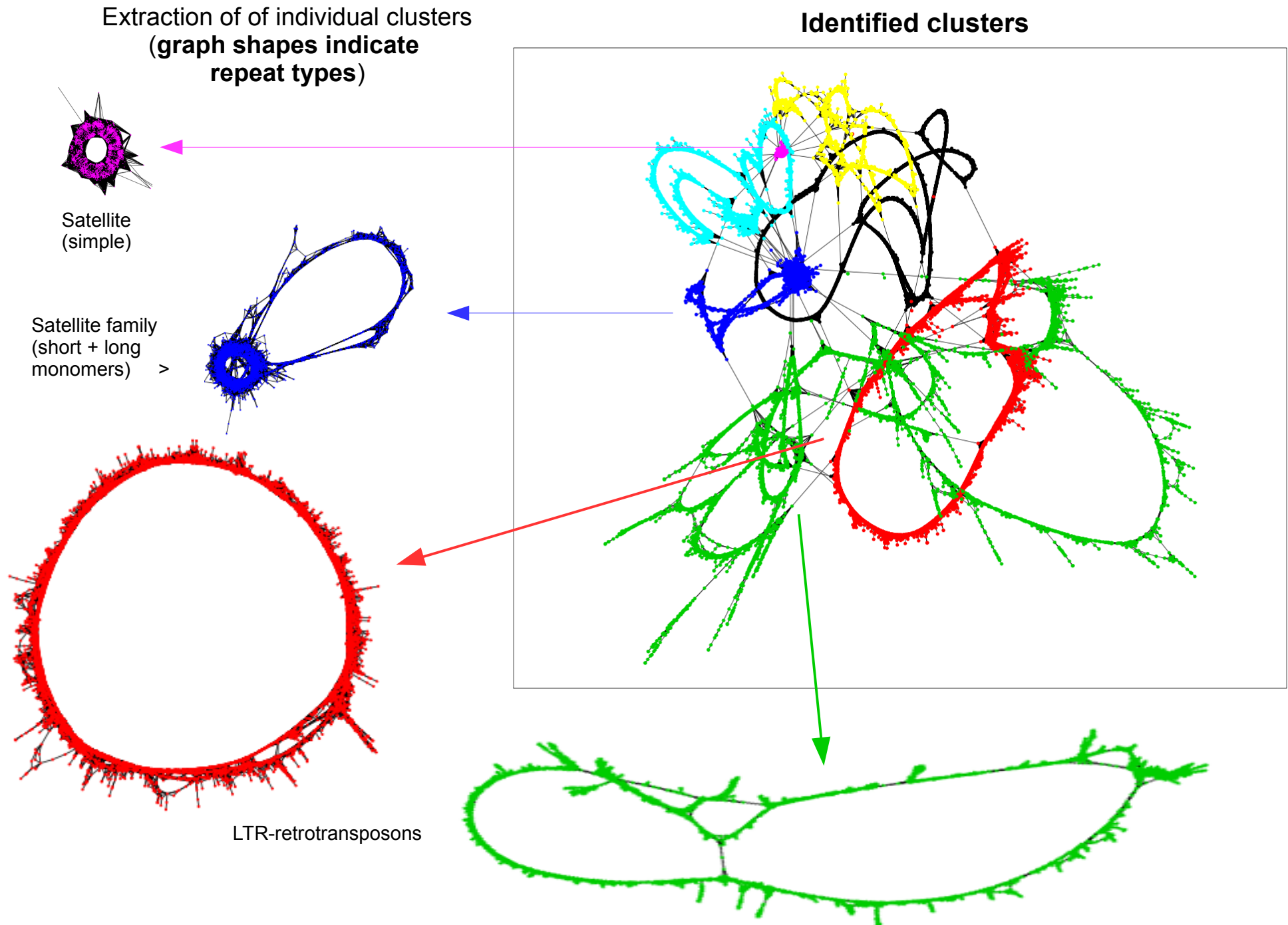
Graph-based characterization of repeat clusters



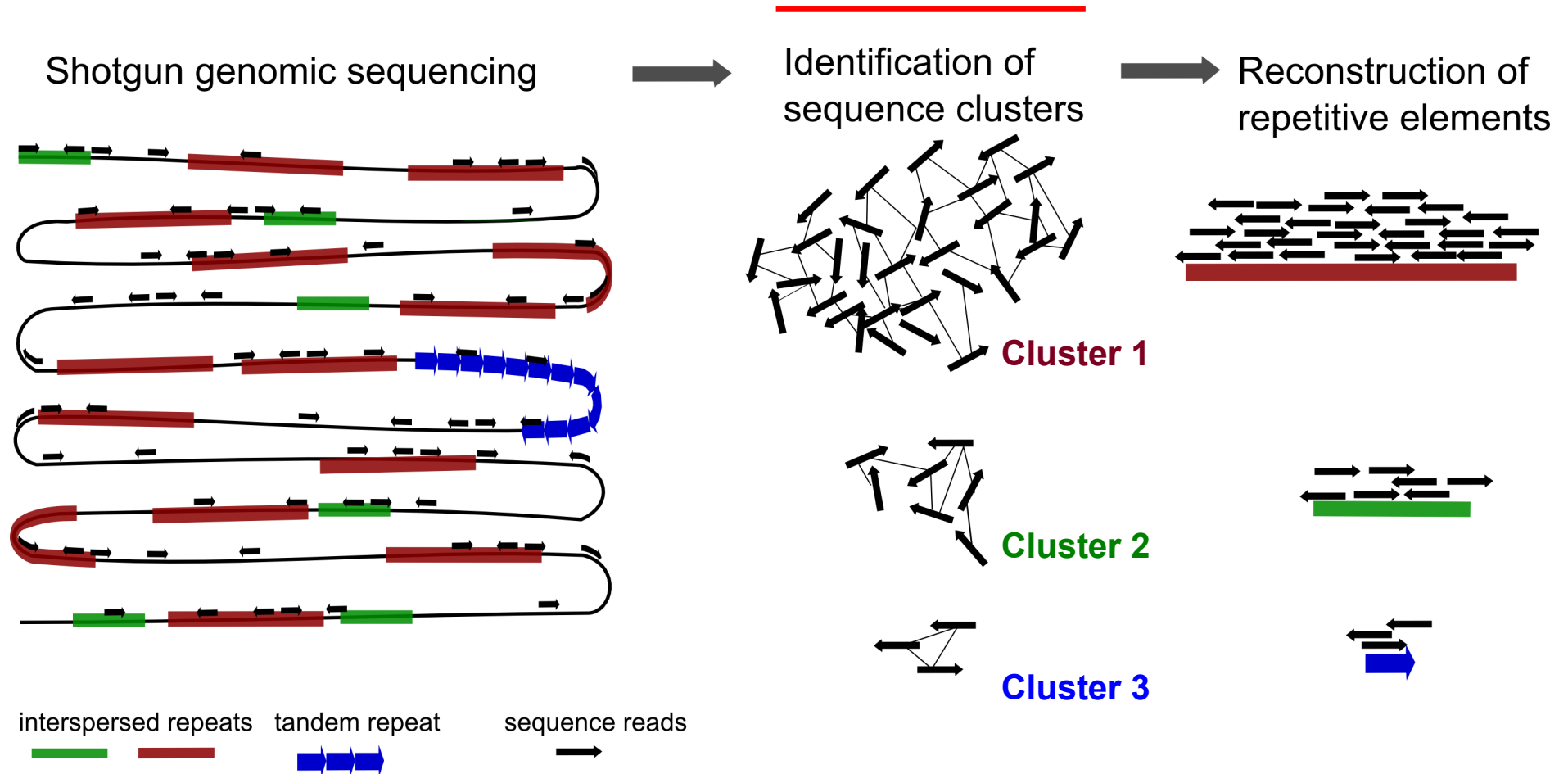
Graph-based clustering

- Sequence overlaps between the reads are transformed to a graph where the **reads** are represented as **nodes** and their **similarities** as **edges** connecting the nodes
- Graph structure is examined to detect *communities of frequently connected nodes* which are split to separate clusters

Graph-based characterization of repeat clusters



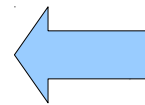
Clustering-based repeat identification



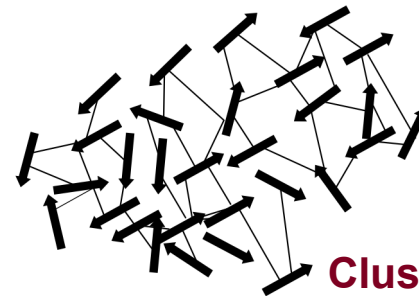
CLUSTER = a set of frequently overlapping reads = REPEAT FAMILY

Repeat characterization

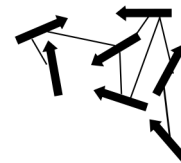
Cluster annotation and quantification



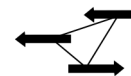
Identification of sequence clusters



Cluster 1



Cluster 2



Cluster 3

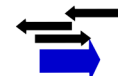
...

...

(Cluster x)



Reconstruction of repetitive elements



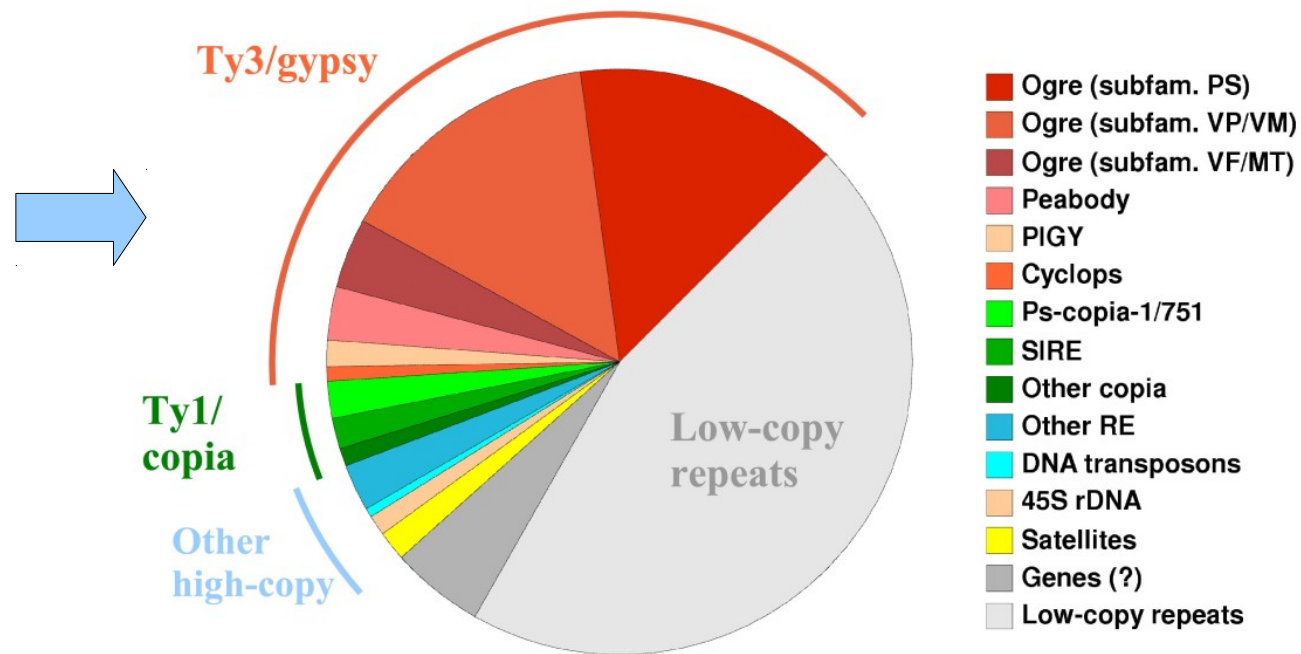
	7192836 (= 100%)	64.1			
CL	reads	genome %	class	type	note
1	304159	4.229	gypsy	Tat	PROT-RT/RH-INT
2	234749	3.264			?, PE->24,18
3	216307	3.007	gypsy	chromo	ALL domains
4	202822	2.820	copla	Maximus	ALL domains
5	149693	2.081	gypsy	Athila	ALL domains
6	145911	2.029	gypsy	Tat	ALL domains
7	143766	1.999	gypsy	chromo	ALL domains
8	142608	1.983	copla	Maximus	ALL domains
9	141836	1.972	LINE		RT
10	123886	1.722	gypsy	chromo	GAG
11	79345	1.103			?,PE->21,95
12	72781	1.012	copla	Angela	ALL domains
13	67096	0.933	gypsy	Tat	ALL domains
14	65455	0.910	gypsy	Athila	GAG, PE->1(!!!),36
15	62334	0.867	gypsy	Tat	ALL domains
16	53845	0.749	copla	Ivana/Oryco	ALL domains
17	49341	0.686			? DNA transp ?? + TR
18	45062	0.626			?,PE->2,63
19	44762	0.622			?,PE->28
20	43332	0.602	tandem		monom ?? ~400 (~1200)
21	42344	0.589	gypsy	chromo	ALL domains
22	40125	0.558	gypsy	Tat	PE->15
23	39923	0.555			?,PE->7,73
24	36353	0.505	gypsy	chromo	(GAG), PE->2,3,28
25	35977	0.500			?,PE->2,81
26	35674	0.496			?
27	34829	0.484	rDNA	5S	
28	34534	0.480	gypsy	chromo	PE->29,19,24
29	34302	0.477	gypsy	chromo	PE->28
30	33114	0.460			?
31	32930	0.458			?

Repeat characterization

Cluster annotation and quantification

	7192836 (= 100%)	64.1			
CL	reads	genome %	class	type	
1	304159	4.229	gypsy	Tat	
2	234749	3.264			
3	216307	3.007	gypsy	chromo	
4	202822	2.820	copia	Maximus	
5	149693	2.081	gypsy	Athila	
6	145911	2.029	gypsy	Tat	
7	143766	1.999	gypsy	chromo	
8	142608	1.983	copia	Maximus	
9	141836	1.972	LINE		
10	123886	1.722	gypsy	chromo	
11	79345	1.103			
12	72781	1.012	copia	Angela	
13	67096	0.933	gypsy	Tat	
14	65455	0.910	gypsy	Athila	
15	62334	0.867	gypsy	Tat	
16	53845	0.749	copia	Ivana/Oryco	
17	49341	0.686			
18	45062	0.626			
19	44762	0.622			
20	43332	0.602	tandem		
21	42344	0.589	gypsy	chromo	
22	40125	0.558	gypsy	Tat	
23	39923	0.555			
24	36353	0.505	gypsy	chromo	
25	35977	0.500			
26	35674	0.496			
27	34829	0.484	rDNA	5S	
28	34534	0.480	gypsy	chromo	
29	34302	0.477	gypsy	chromo	
30	33114	0.460			
31	32930	0.458			

Proportions of various repeat types in a genome



RepeatExplorer

- *What RepeatExplorer can do for you:*

- Identify all sequences with certain number of repetitions
- Repeat quantification
- Provide models of repeat populations (sequence variability)
- Help in repeat classification/annotation (in plant genomes, *Metazoa in prep.*)

- *What it cannot do:*

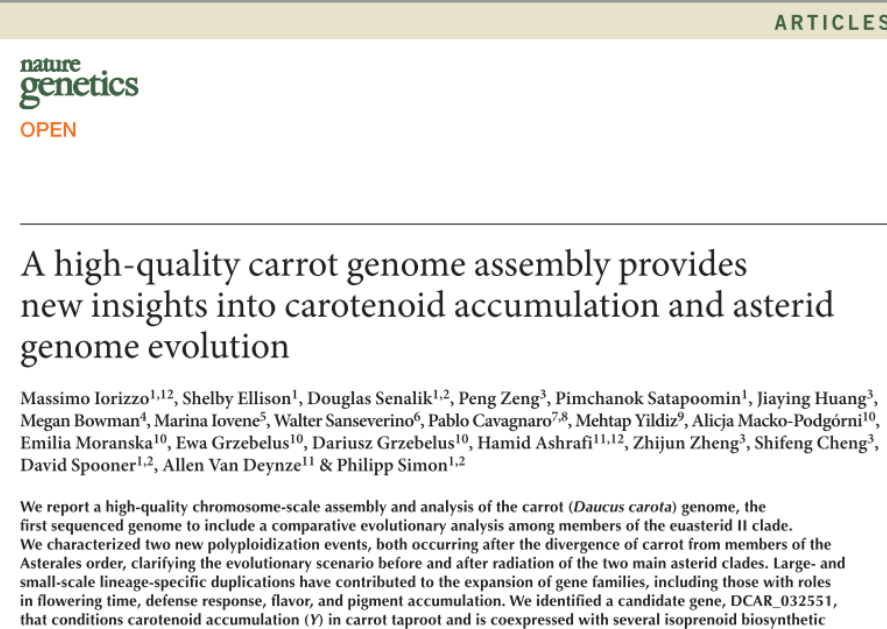
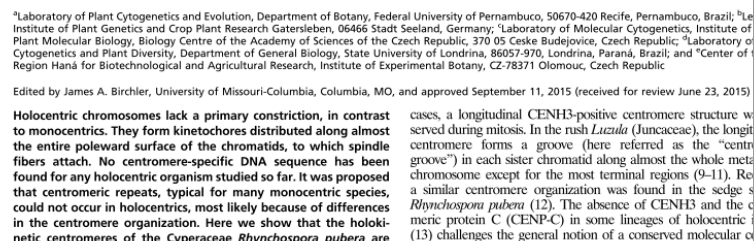
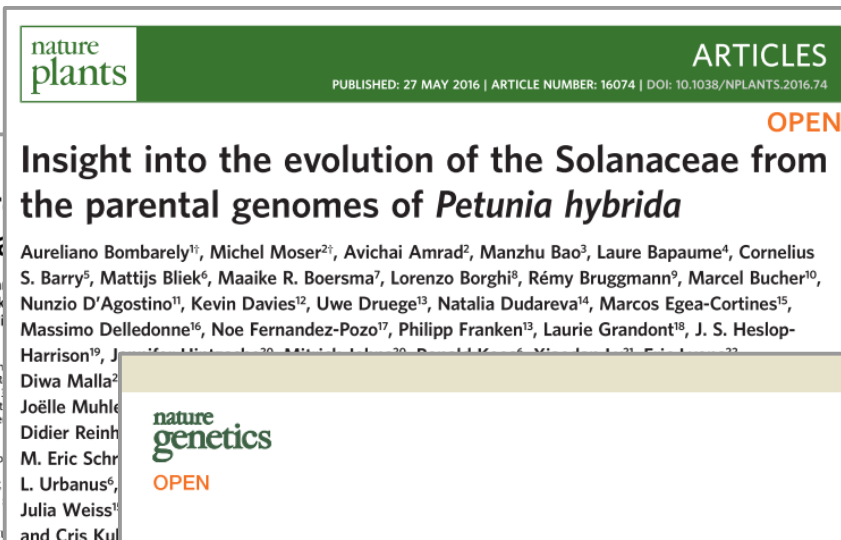
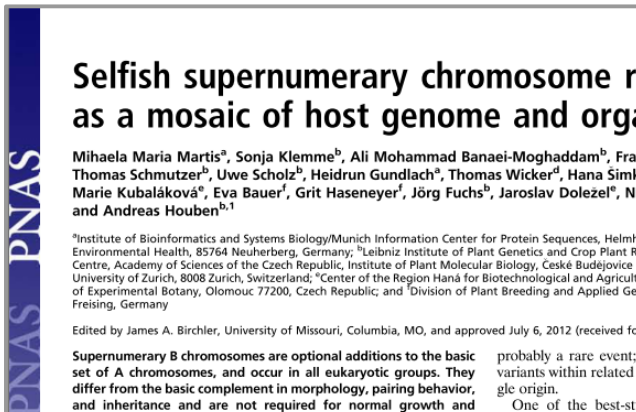
- Genome assembly or reconstruction of individual repeat copies
- Analyze repeats using assembled genomes as input
- Use long NGS reads (PacBio, Oxford Nanopore) as input (*work in progress*)
- Identify some tandem repeats with very short monomers or other low-complexity repeats

Think about proper design of your analysis - RE is just a program designed for a specific type of input (e.g. it needs WGS data as input and it will not work with BAC clones)

Repetitive DNA characterization using RepeatExplorer

Plants

- Over 100 species characterized so far
- Comparative studies
- Whole genome assembly projects



Repetitive DNA characterization using RepeatExplorer

Plants

- Over 100 species characterized so far
- Comparative studies
- Whole genome assembly projects

Mammals

Bats, deer

Fish

- Austrolebias charrua*, *Cynopoecilus melanotaenia*

Insects

- Locust, grasshoppers, kissing bugs

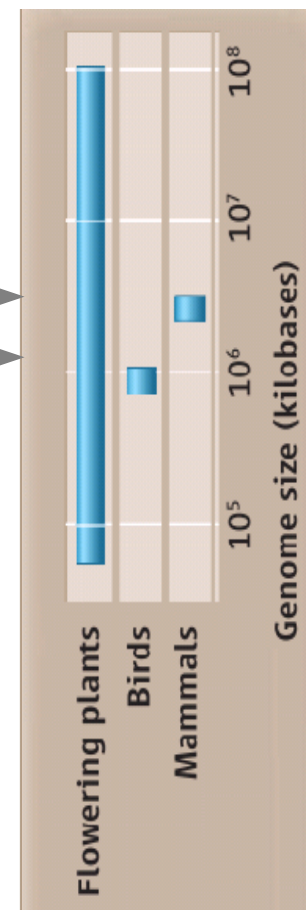
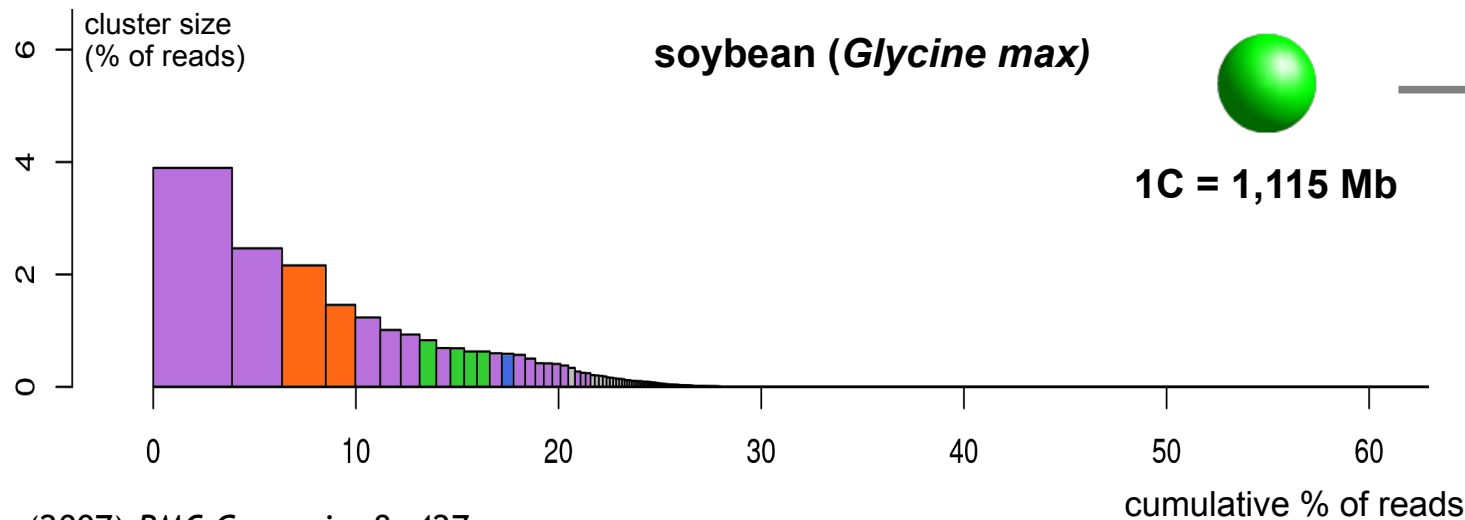
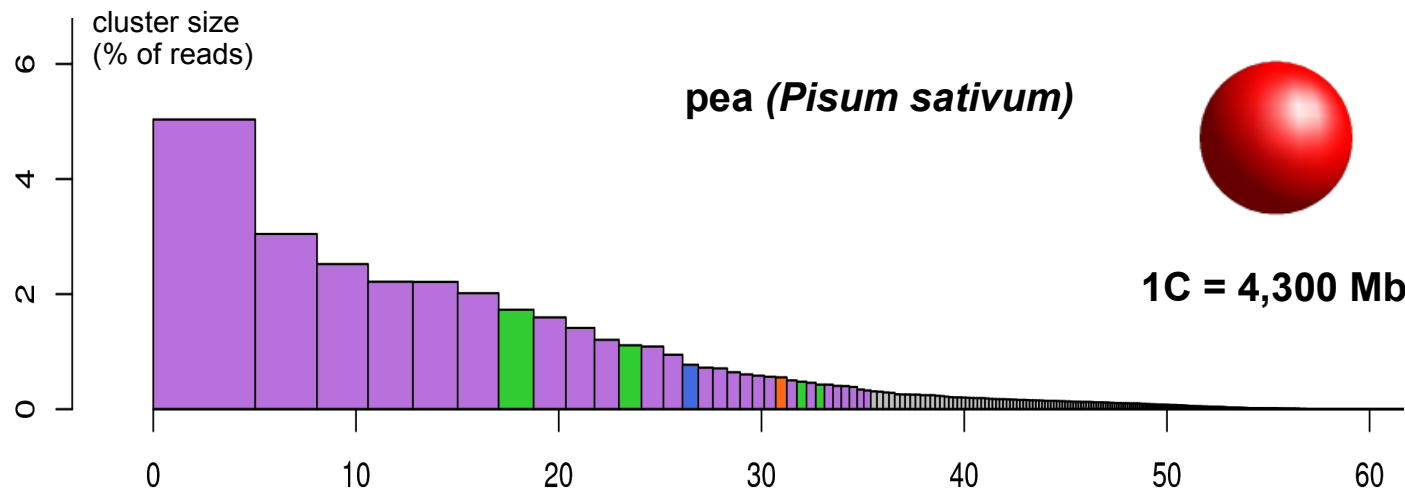
Worms

- Soil helminths

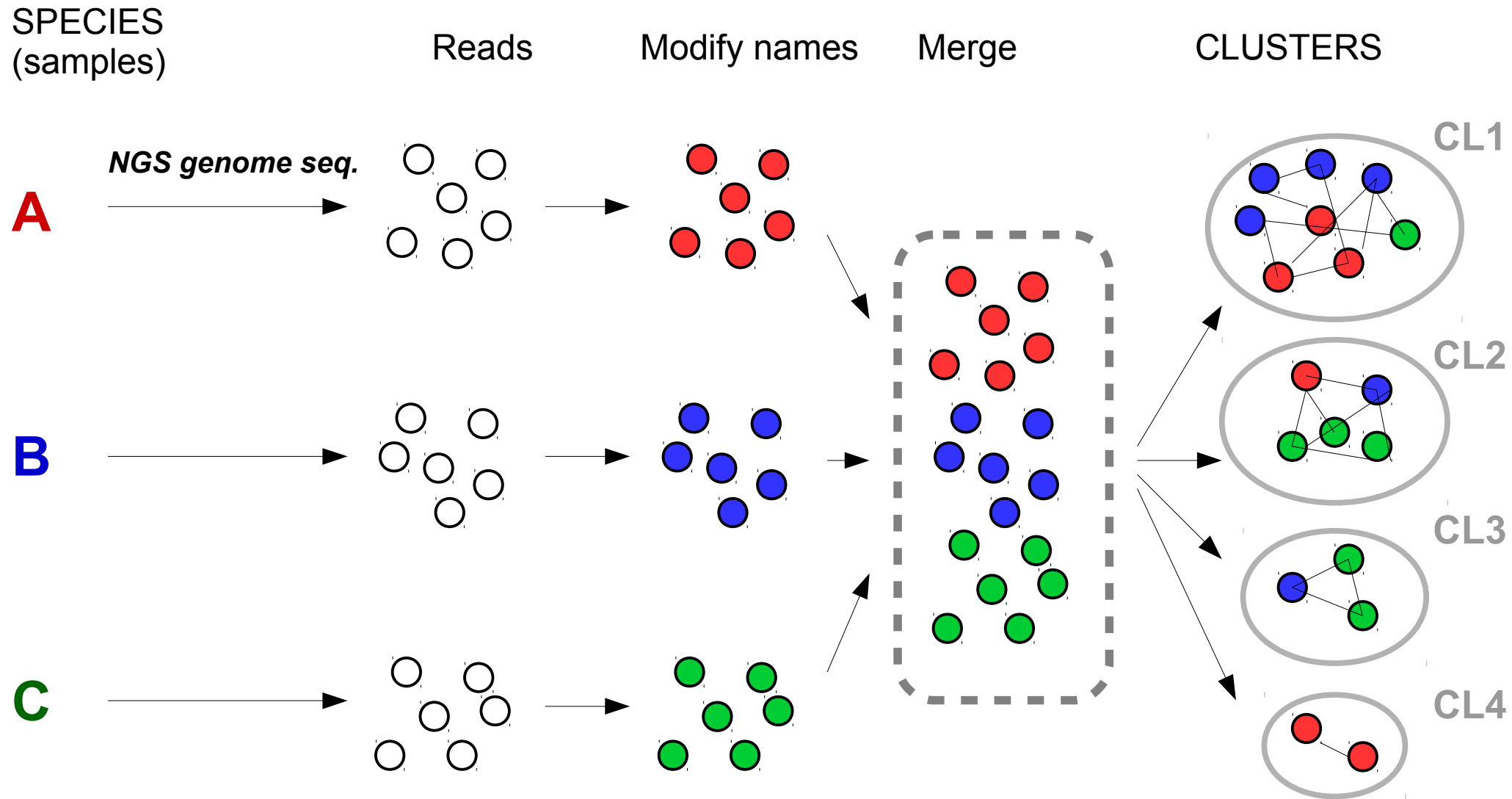


APPLICATIONS: Repeat composition of individual species

Repeat type: Ty3/gypsy Ty1/copia Satellite DNA rDNA other/unknown



APPLICATIONS: Comparative analysis



Comparative analysis

Sequence reads from **multiple species are mixed** and subjected to clustering analysis

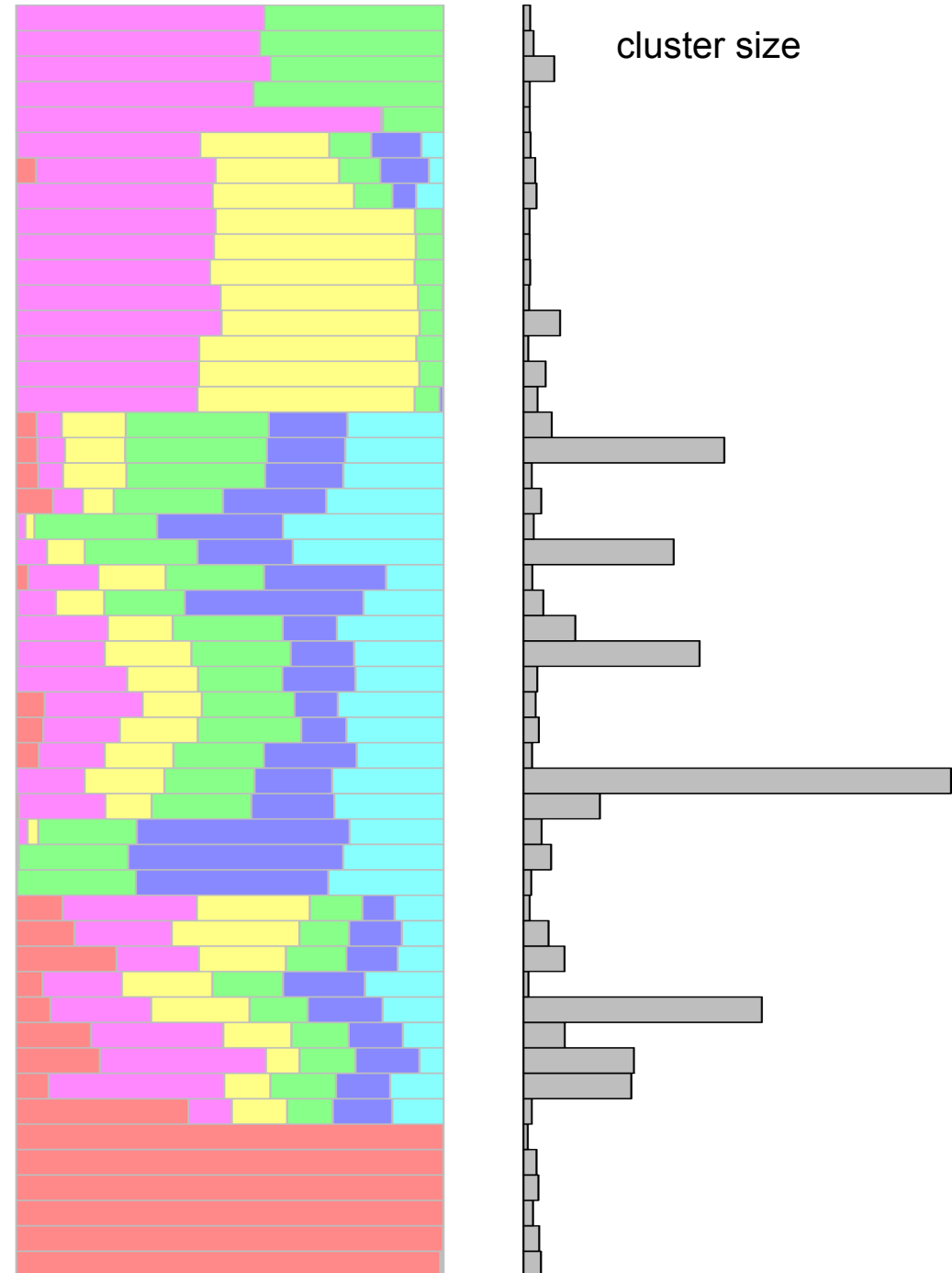
- Efficient identification of homologous repeats from different species
- Easy quantification



SPECIES:

	<i>Musa ornata</i>
	<i>M. balbisiana</i> (B)
	<i>M. acuminata</i> (A)
	<i>M. textillis</i> (T)
	<i>M. beccarii</i>
	<i>Ensete gillettii</i>

proportion
of reads from
individual
species >>>



Comparative analysis

A

SPECIES:

Ensete gillettii

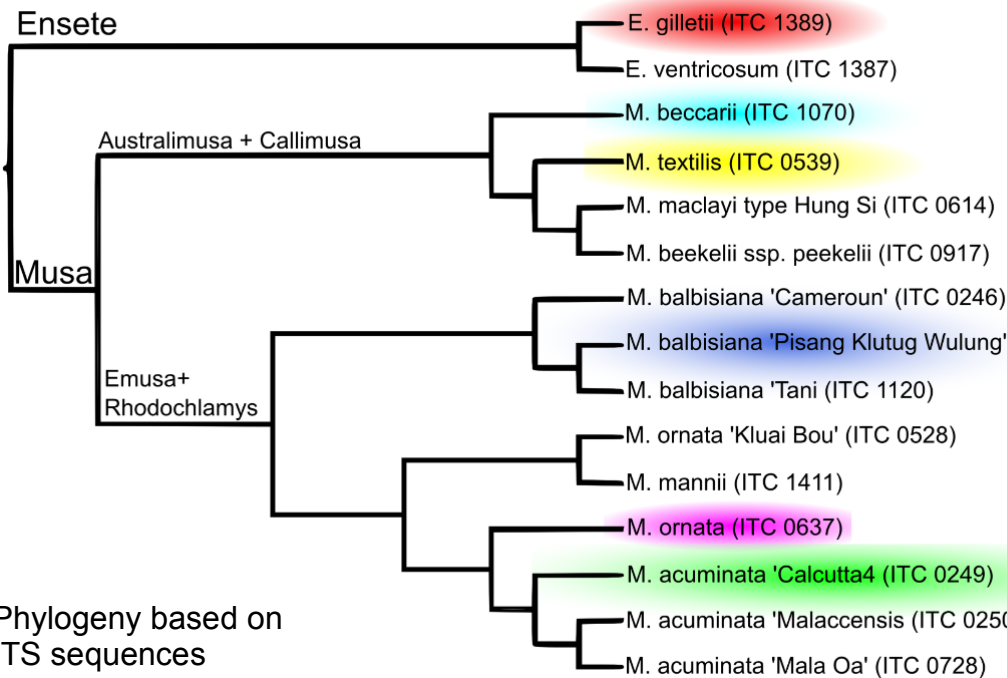
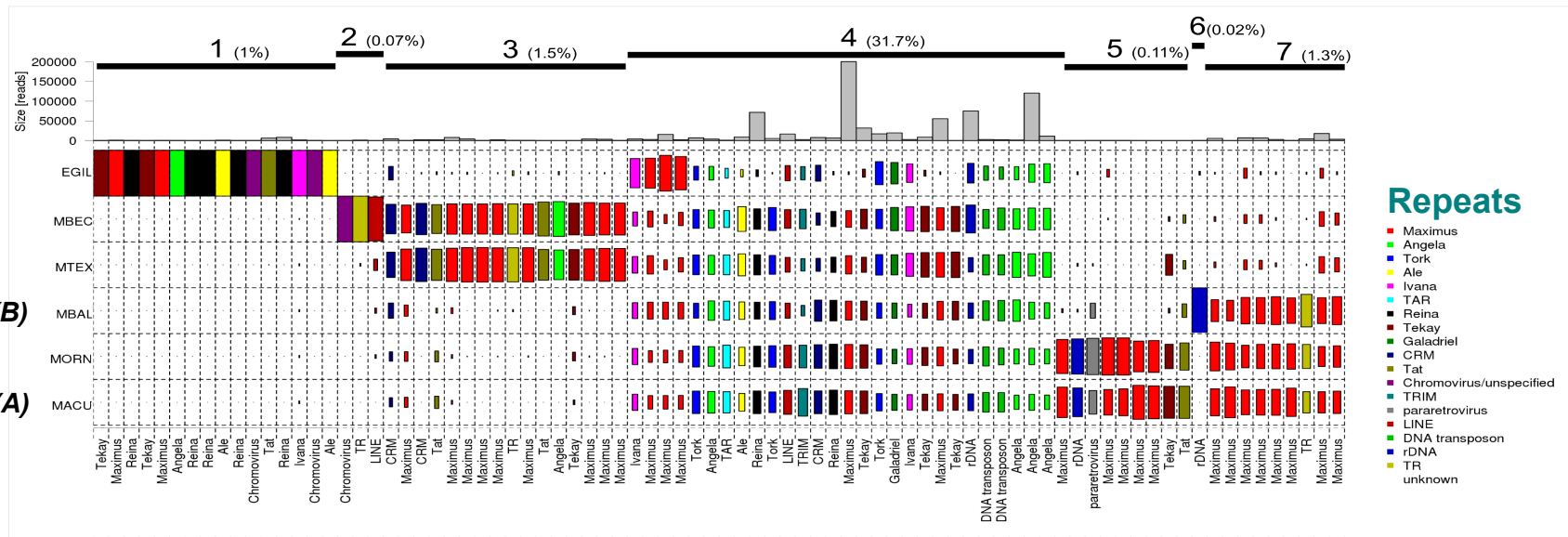
Musa beccarii

M. textillis (T)

M. balbisiana (B)

M. ornata

M. acuminata (A)



Phylogeny based on
ITS sequences



Comparative analysis

A

SPECIES:

Ensete gillettii

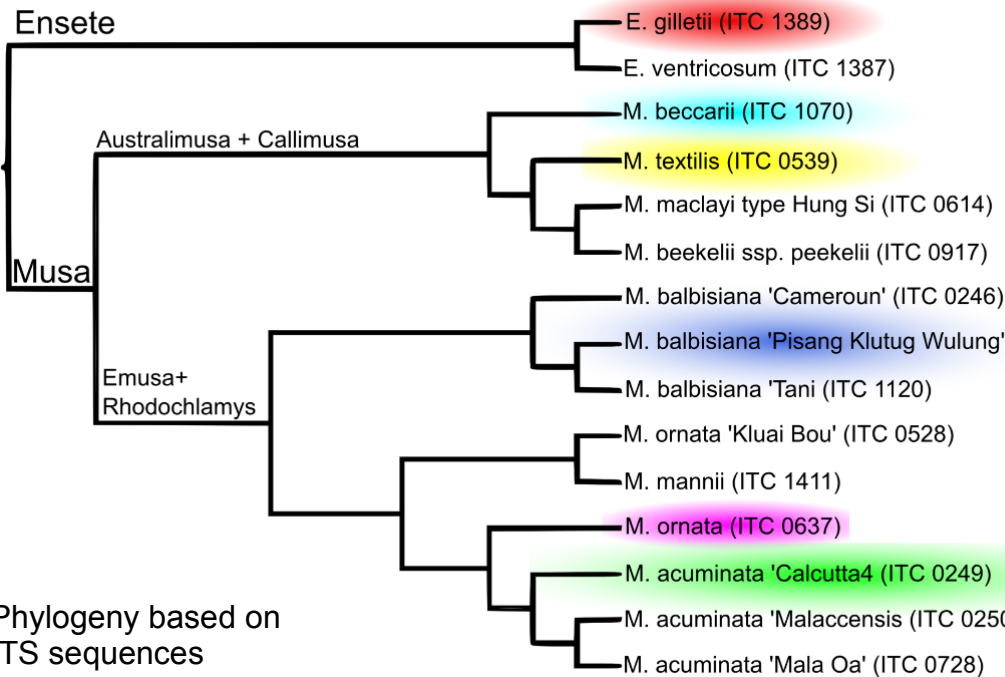
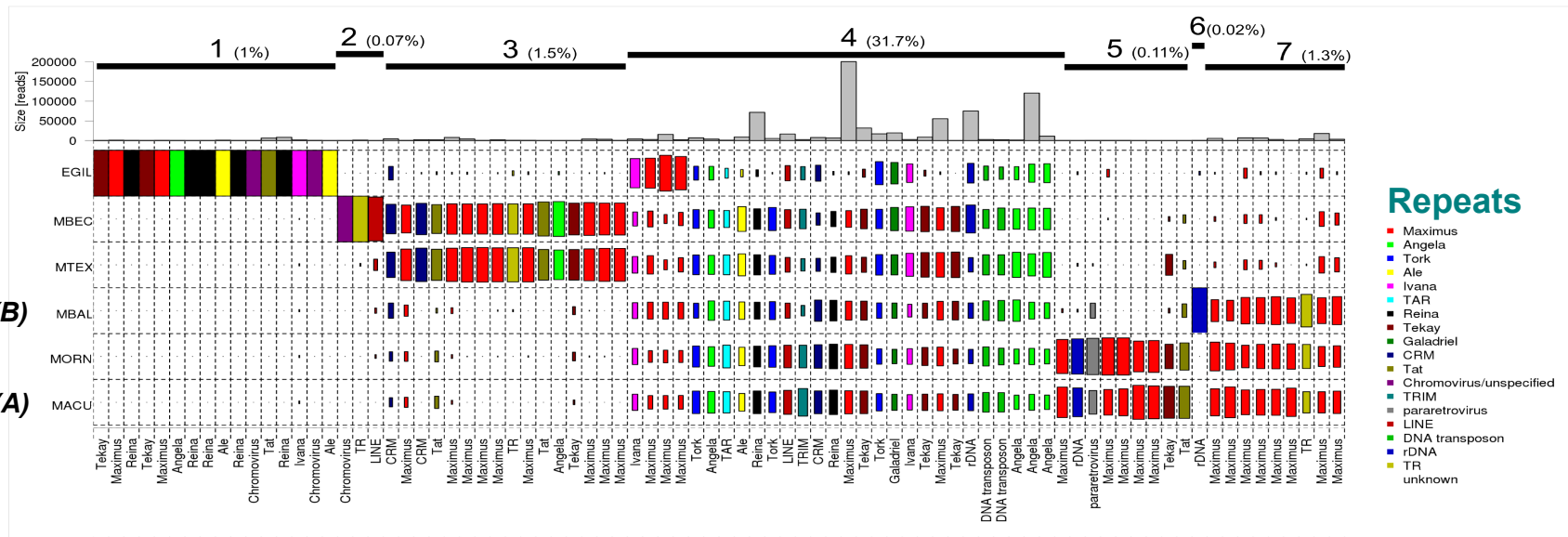
Musa beccarii

M. textillis (T)

M. balbisiana (B)

M. ornata

M. acuminata (A)



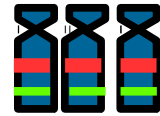
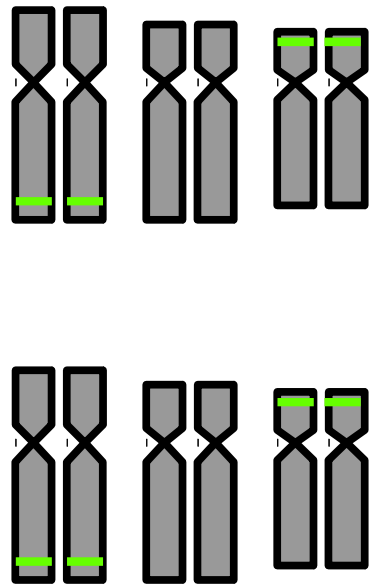
Phylogeny based on
ITS sequences

Genomic repeat abundances
contain phylogenetic signal

Dodsworth et al. (2015)
Syst. Biol. 64(1): 112-126



Comparative analysis (searching for B-specific repeats)



Next-gen
sequencing

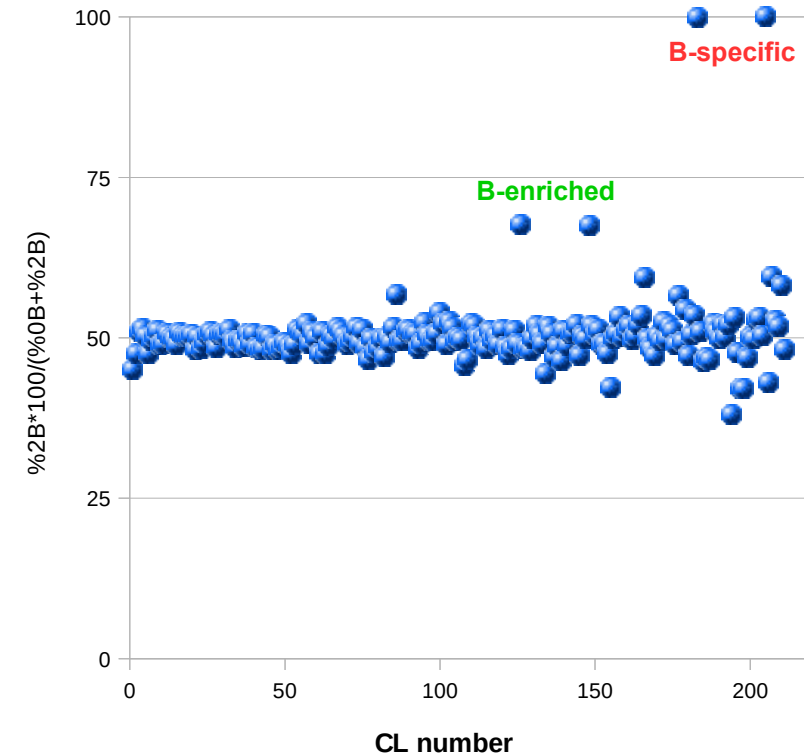
As + Bs

Next-gen
sequencing

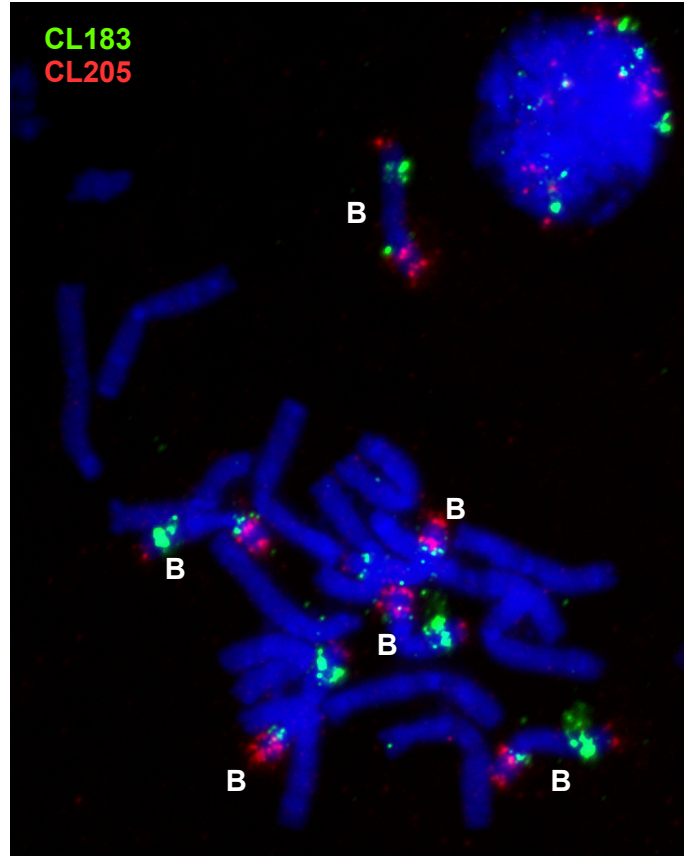
As - 0 -

Mix coded
reads
and
perform
CLUSTERING

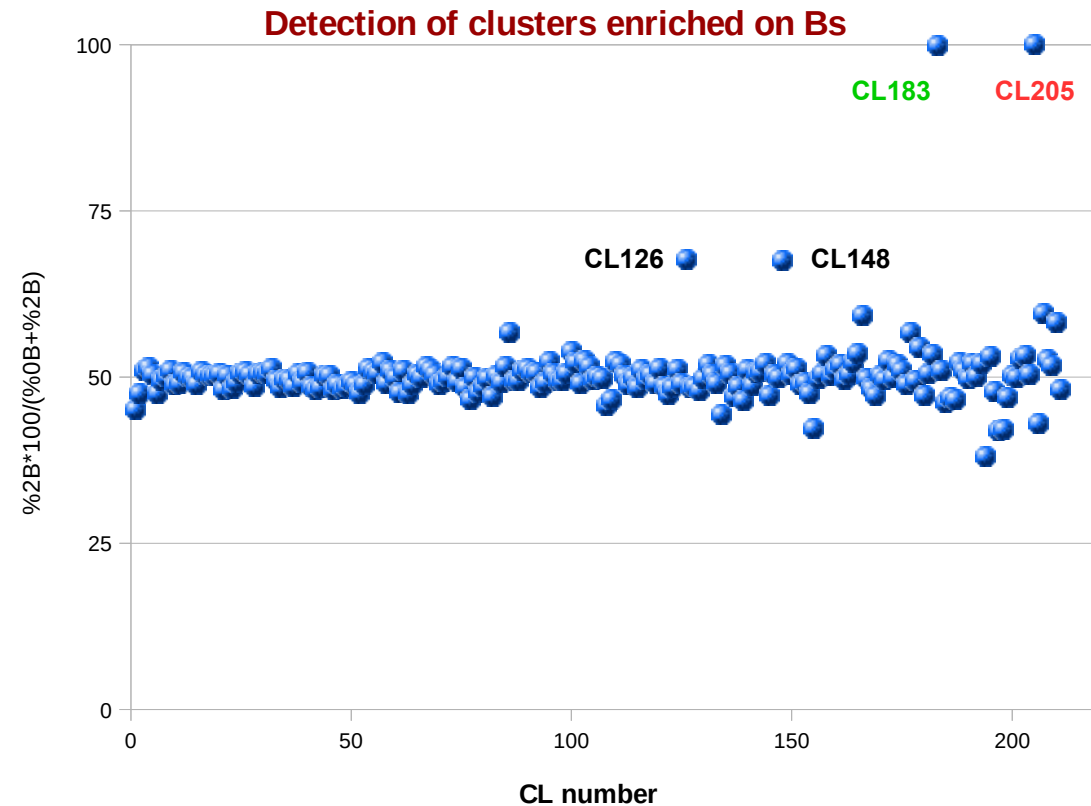
Detection of clusters enriched on Bs



Aegilops speltoides, comparative analysis of 0B / 2B plants



FISH by Alevtina Ruban



Wu et al. (2019) *New Phytol.*, in press

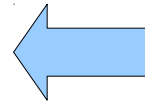
Repeats significantly enriched on Bs

CL	reads	[%]	0B	%	2B	%	$\frac{\%2B \cdot 100}{\%2B + \%0B}$	Repeat type
126	5216	0.191	1688	0.120	3528	0.252	68	satellite (86 bp)
148	2877	0.105	934	0.067	1943	0.139	68	tandem (?)
183	731	0.027	1	0.000	730	0.052	100	satellite (~1.1 kb)
205	358	0.013	0	0.000	358	0.026	100	satellite (185 bp)

Cluster-centered downstream applications

Cluster annotation and quantification

CL	reads	genome %	class	type
1	304159	4.229	gypsy	Tat
2	234749	3.264		
3	216307	3.007	gypsy	chromo
4	202822	2.820	copia	Maximus
5	149693	2.081	gypsy	Athila
6	145911	2.029	gypsy	Tat
7	143766	1.999	gypsy	chromo
8	142608	1.983	copia	Maximus
9	141836	1.972	LINE	
10	123886	1.722	gypsy	chromo
11	79345	1.103		
12	72781	1.012	copia	Angela
13	67096	0.933	gypsy	Tat
14	65455	0.910	gypsy	Athila
15	62334	0.867	gypsy	Tat
16	53845	0.749	copia	Ivana/Oryco
17	49341	0.686		
18	45062	0.626		
19	44762	0.622		
20	43332	0.602	tandem	
21	42344	0.589	gypsy	chromo
22	40125	0.558	gypsy	Tat
23	39923	0.555		
24	36353	0.505	gypsy	chromo
25	35977	0.500		
26	35674	0.496		
27	34829	0.484	rDNA	5S
28	34534	0.480	gypsy	chromo
29	34302	0.477	gypsy	chromo
30	33114	0.460		
31	32930	0.458		



CLUSTERS

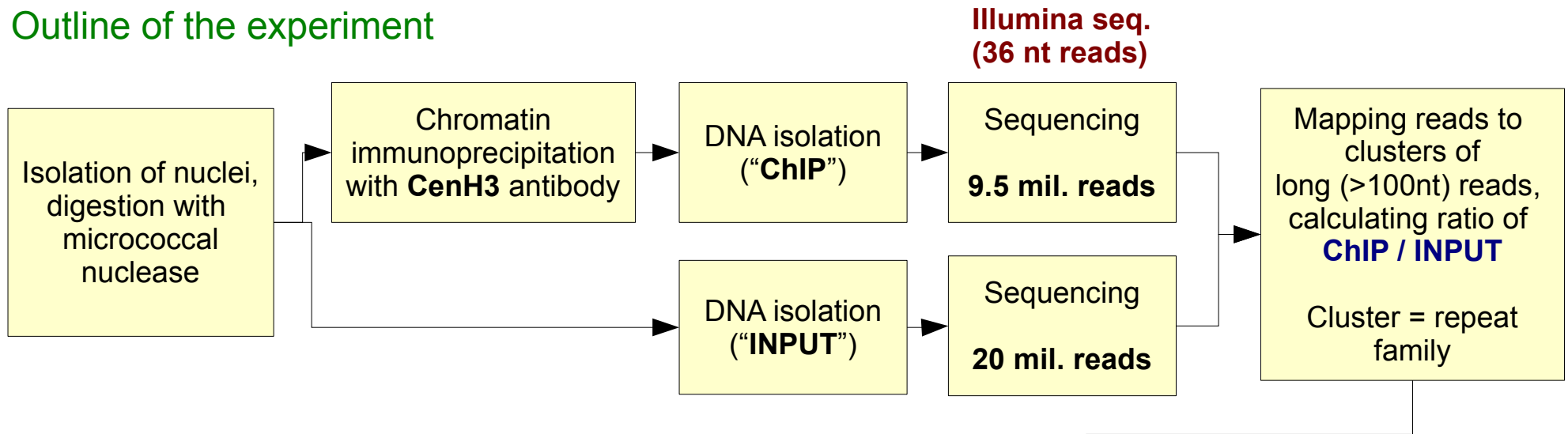
- Represent specific repeat families/variants or their parts
- They are collections of sequence reads capturing full sequence variability of repeat populations

Using clusters as reference for similarity-based classification of:

- RNA-seq reads (mRNA, smRNAs,...)
 - Detection of transcribed repeats
 - Comparative analysis (tissues,...)
- ChIP-seq reads
 - Association of repeats with specific types of chromatin

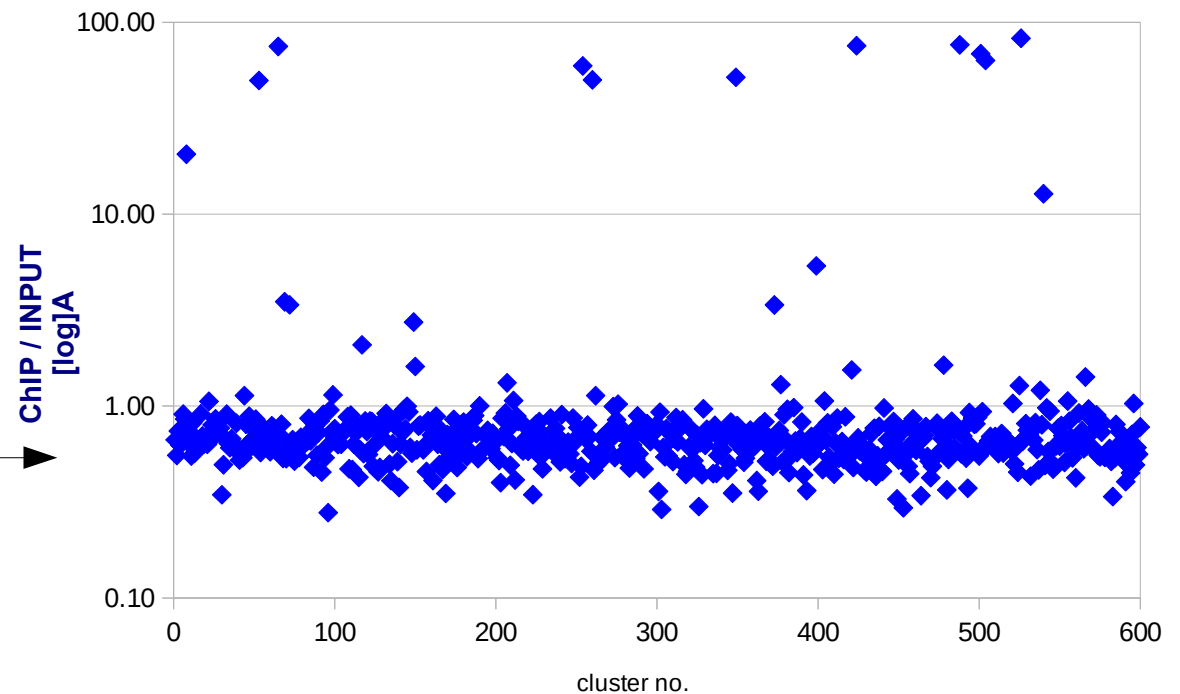
Identification of centromeric repeats by ChIP-seq

Outline of the experiment



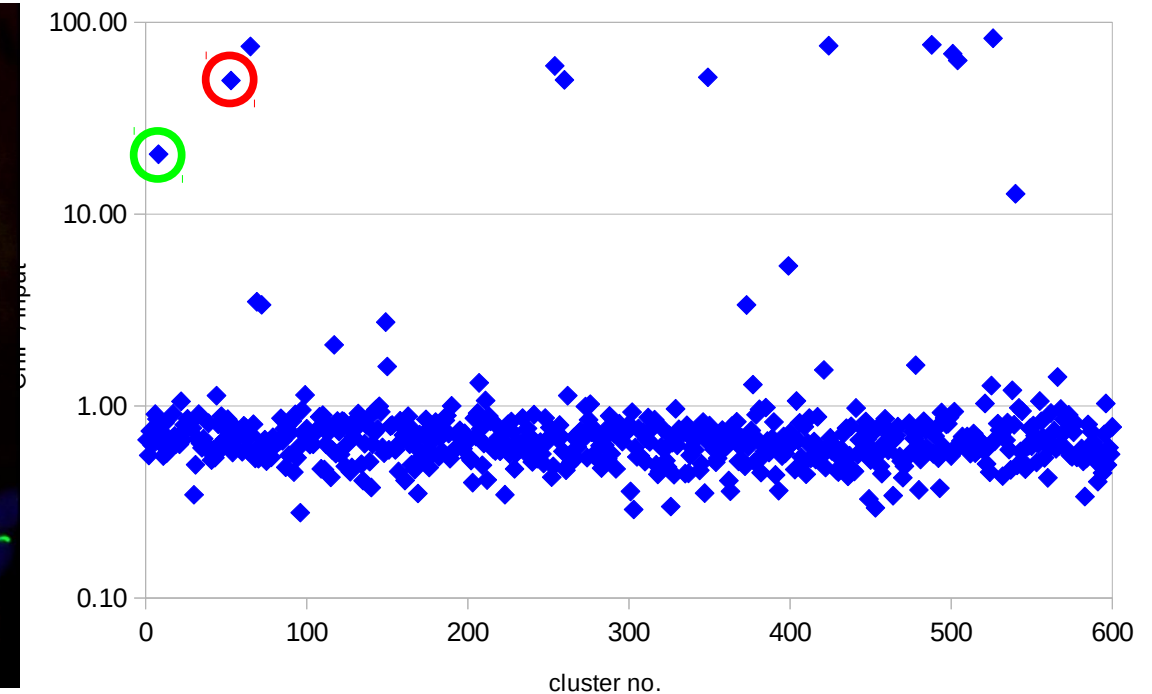
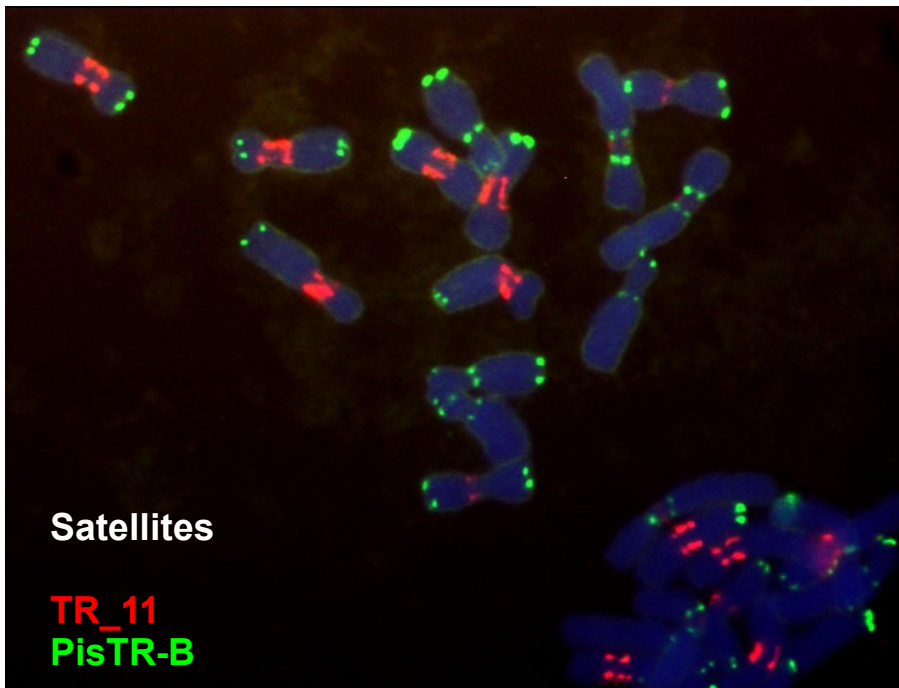
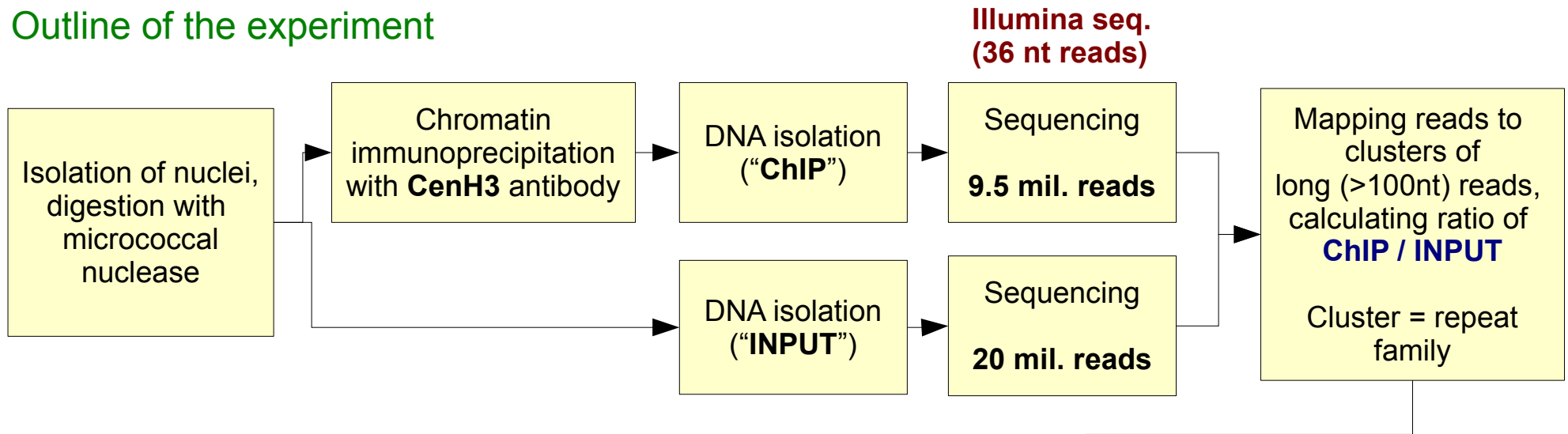
100-fold enrichment

no enrichment



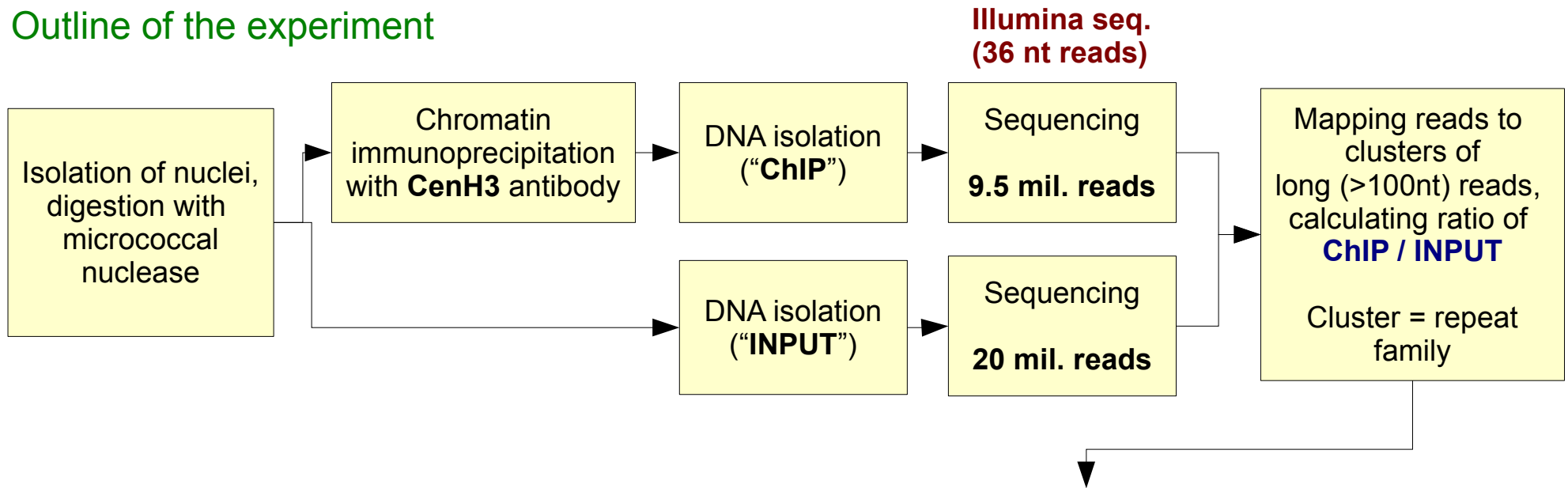
Identification of centromeric repeats by ChIP-seq

Outline of the experiment



Identification of centromeric repeats by ChIP-seq

Outline of the experiment



ChIP-seq Mapper

now available as a tool
on our Galaxy server

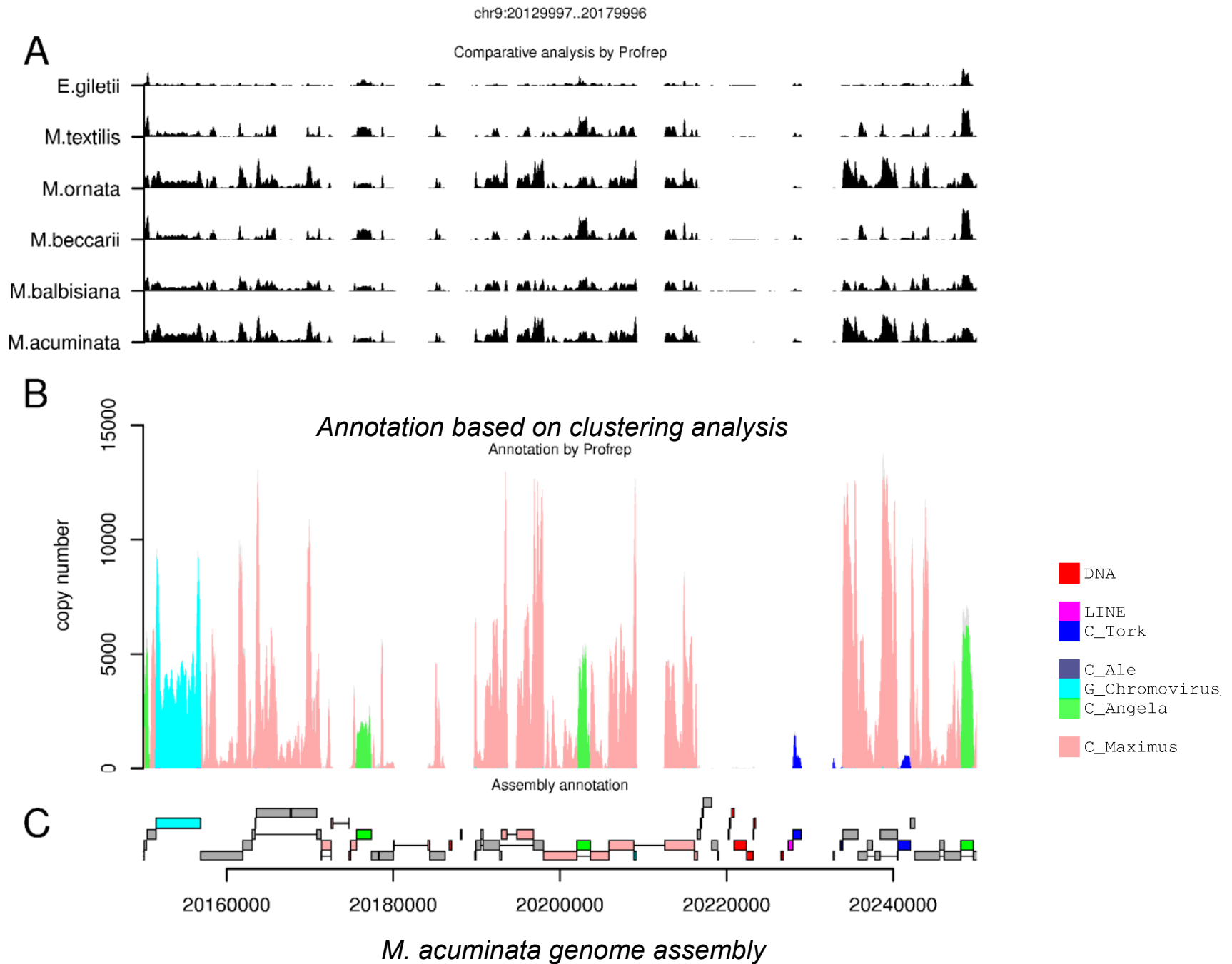
>>>

The screenshot shows the Galaxy web interface with the ChIP-seq Mapper tool selected. The tool is titled "Chip-Seq Mapper (Galaxy Version 0.1.1)". It has a sidebar with "Tools" and "Workflows" sections. The main panel contains the following fields:

- Chip Sequences:** A dropdown menu showing "36: PFCC" with a file icon, a copy icon, and a folder icon. Below it, it says "NGS data in fasta format".
- Input Sequences:** A dropdown menu showing "35: PFCI" with a file icon, a copy icon, and a folder icon. Below it, it says "NGS data in fasta format".
- Reference - Contig Sequences:** A dropdown menu showing "34: PFL_max contigs.info" with a file icon, a copy icon, and a folder icon. Below it, it says "Contigs obtained from RepeatExplorer clustering pipeline in fasta file".
- Number of clusters to be shown in graph:** A text input field with the value "600".
- Minimum bit score threshold:** A text input field with the value "30".
- Execute:** A blue button with a checkmark and the text "Execute".

Below the input fields, there is a note: "All similarity hits with lower bit score will not be considered for ChIP/Input ratio calculation".

PROFREP - improving repeat annotation in genome assembly



PROFREP and DANTE

The screenshot shows the Galaxy web interface. On the left, the 'Tools' panel lists 'Experimental Tools' including 'PROFREP DATA PREPARATION' with sub-items 'Extract Data For ProfRep' and 'ProfRep Reducing', and 'PROFREP' with sub-items 'ProfRep' and 'ProfRep Refiner'. The main panel displays the 'Extract Data For ProfRep Extract data for ProfRep from RepaetExplorer (Galaxy Version 1.0.0)' tool. It has a dropdown menu for 'RepaetExplorer output data archive' set to '27: RepaetExplorer2 - Archive with HTML report from data 22' and an 'Execute' button. Below the tool description, it states 'WHAT IT DOES: This tool will extract all the necessary input data that are needed for ProfRep from RE output archive. Use if the species is not already provided in our internal database (ProfRep drop-down menu -> Choose existing annotation dataset)'.

Set of tools available on the Galaxy server

data import from RE2...

... to final annotation

