

Where are we now?



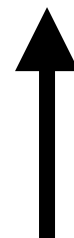
**Exhausting
experiment**



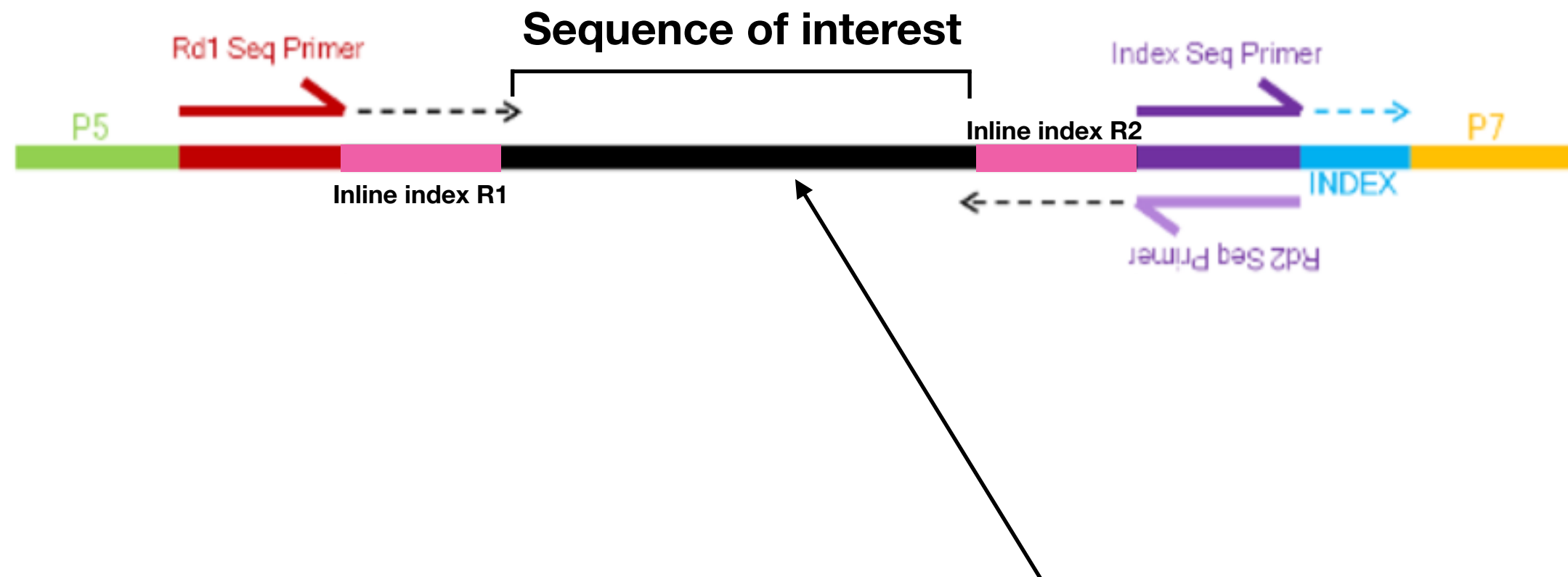
Sequencing

ATCG.....

Analysis our data



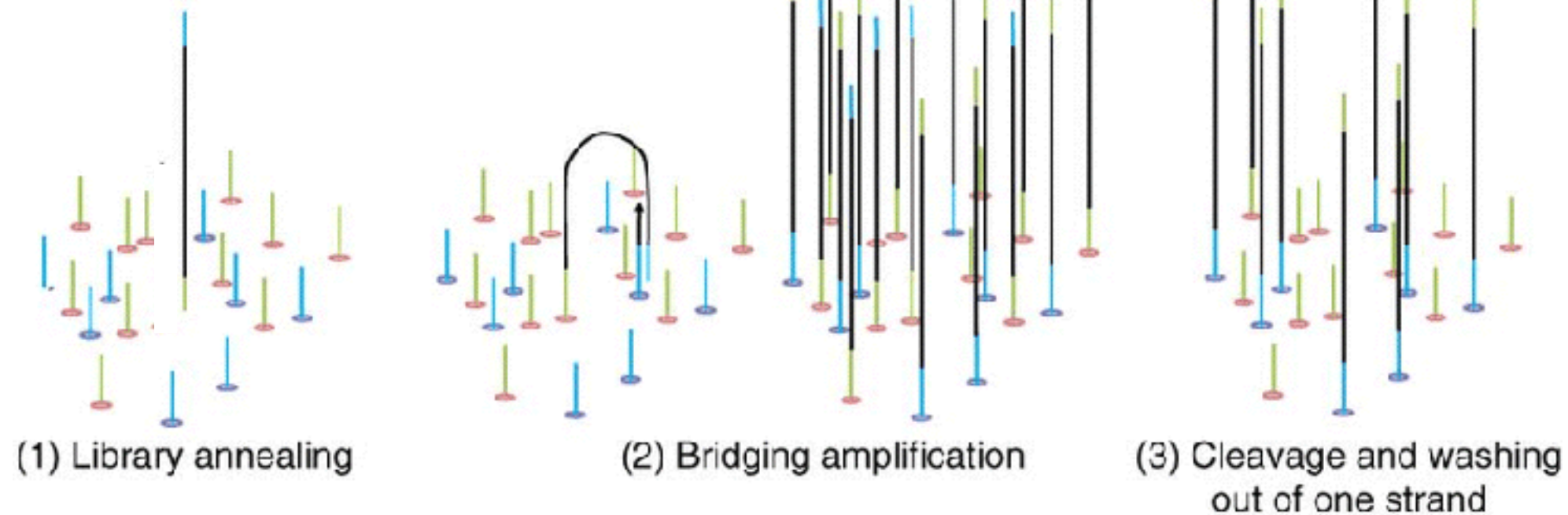
Structure of prepared library



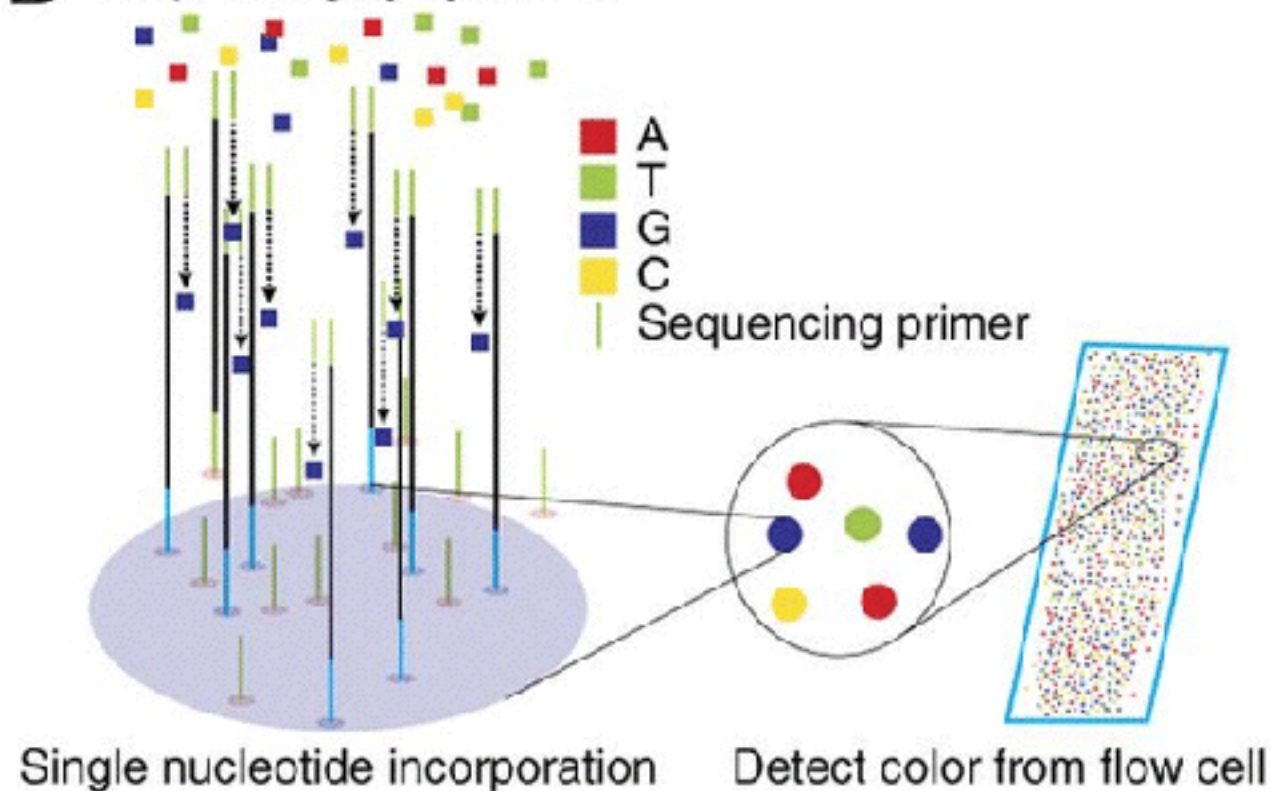
Length of this part is sometimes called “**insert size**”, which totally depends on the length of DNA after shearing

Pair-end Sequencing

A Cluster generation



B Sequencing by synthesis

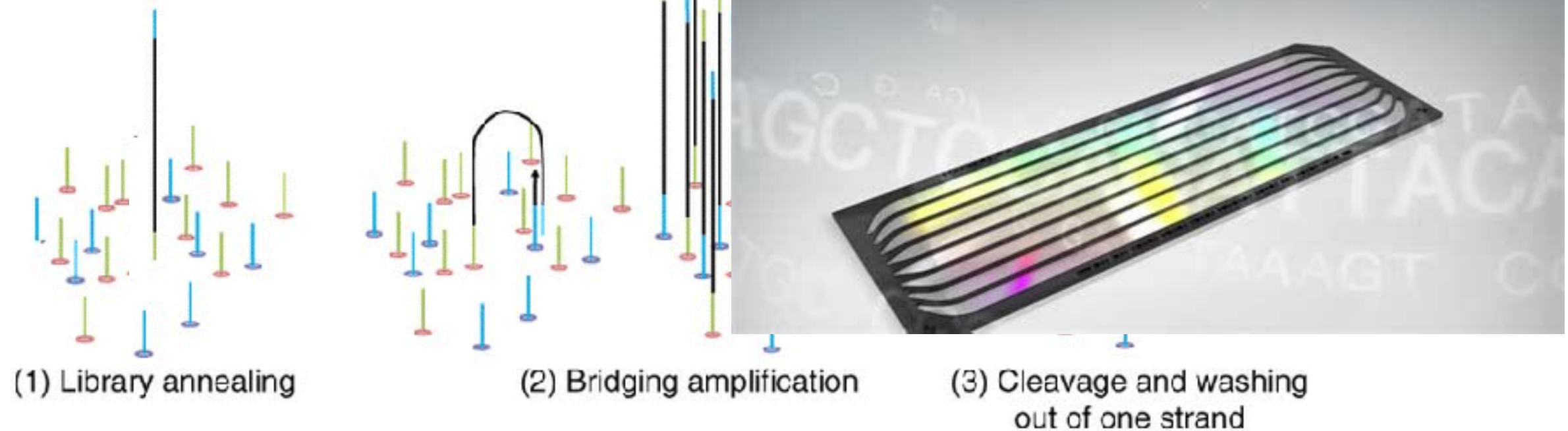


C Preparation for paired-end sequencing

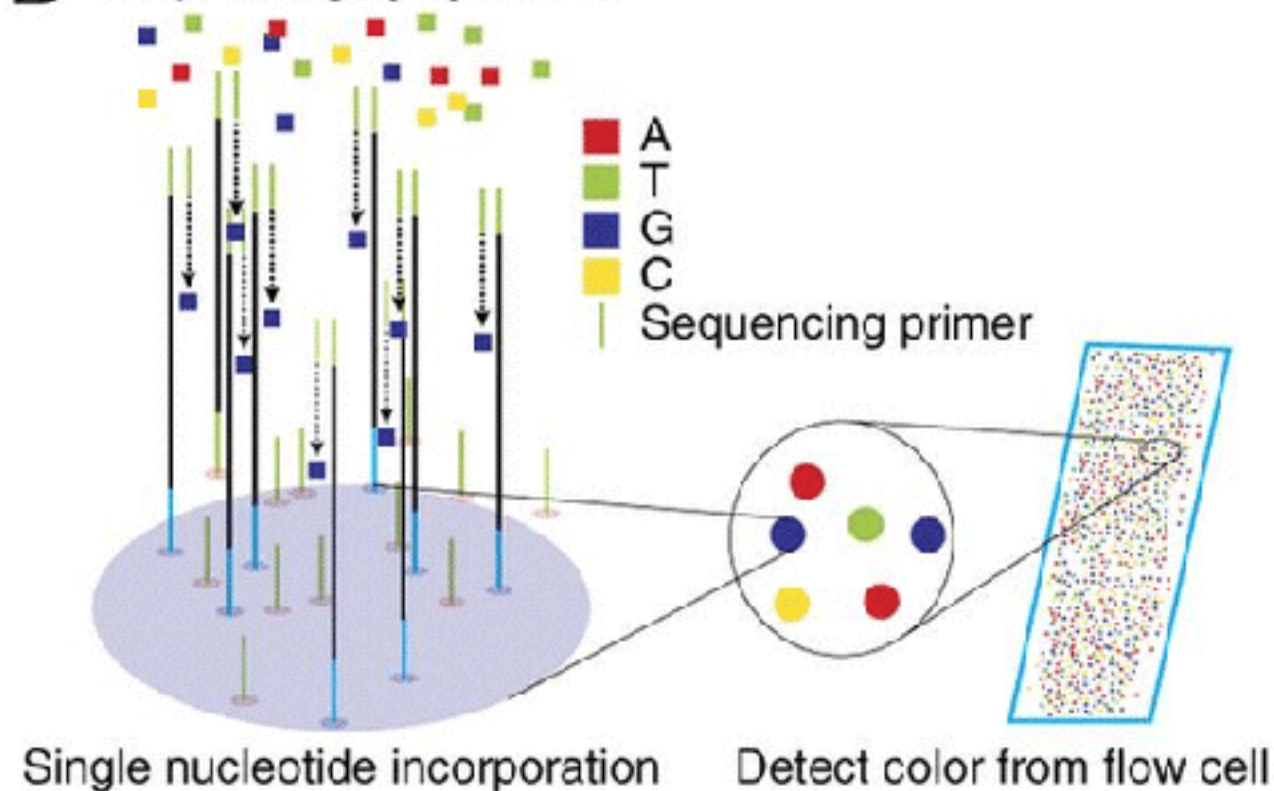


Pair-end Sequencing

A Cluster generation



B Sequencing by synthesis

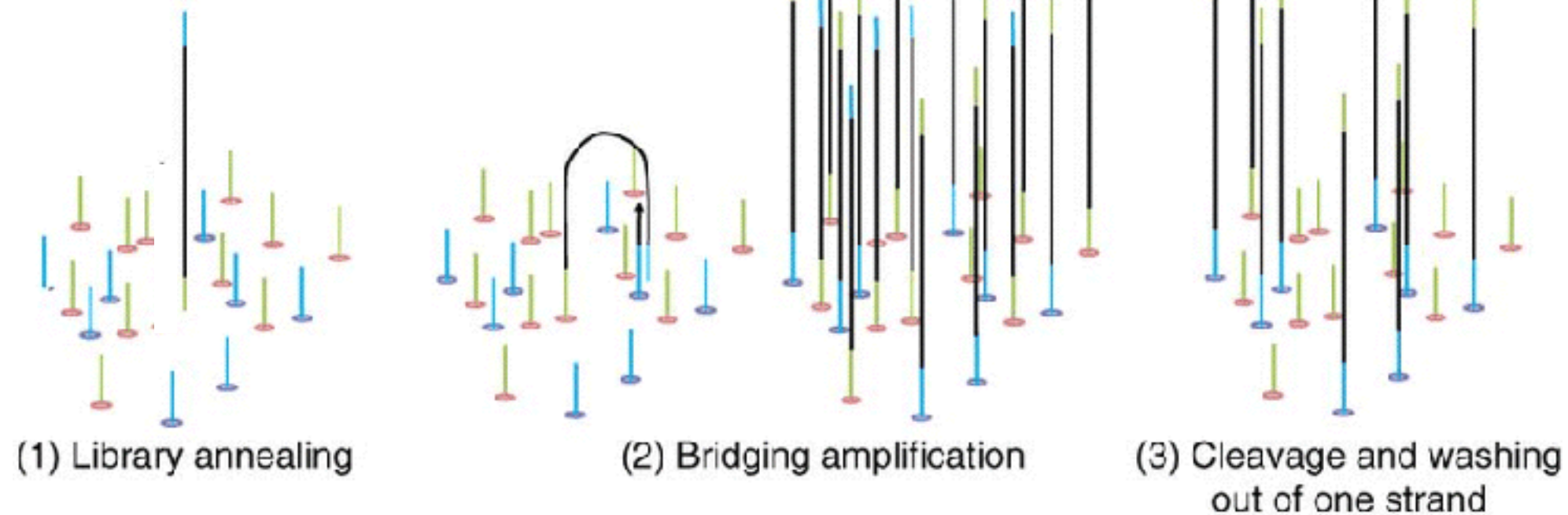


C Preparation for paired-end sequencing

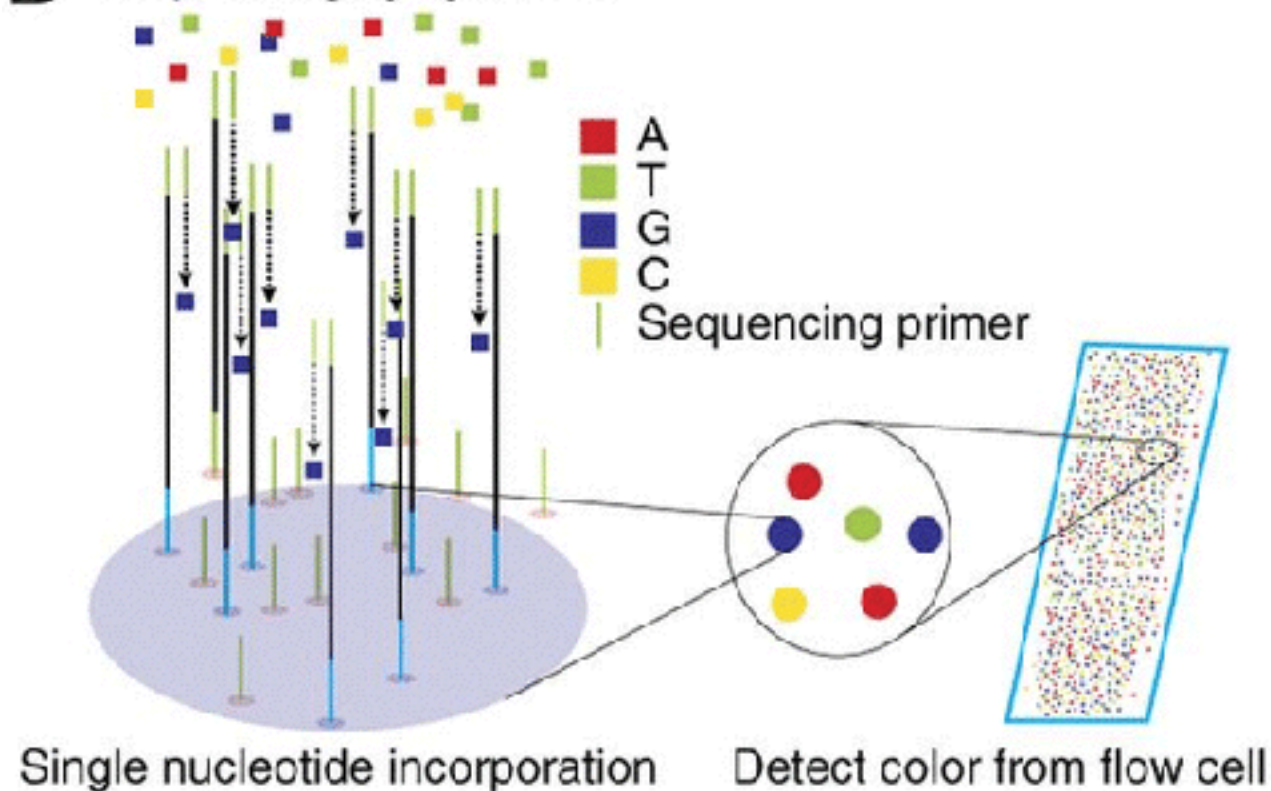


Pair-end Sequencing

A Cluster generation



B Sequencing by synthesis

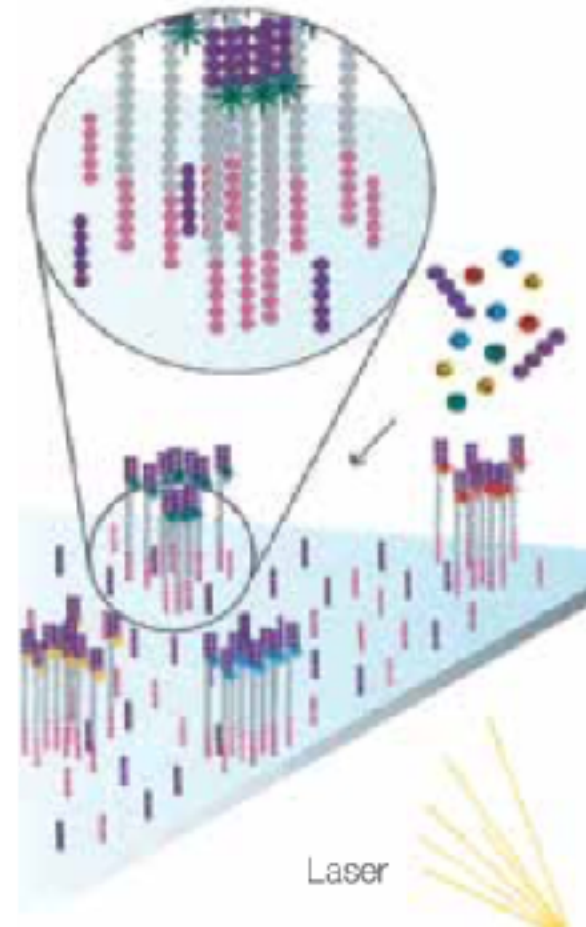
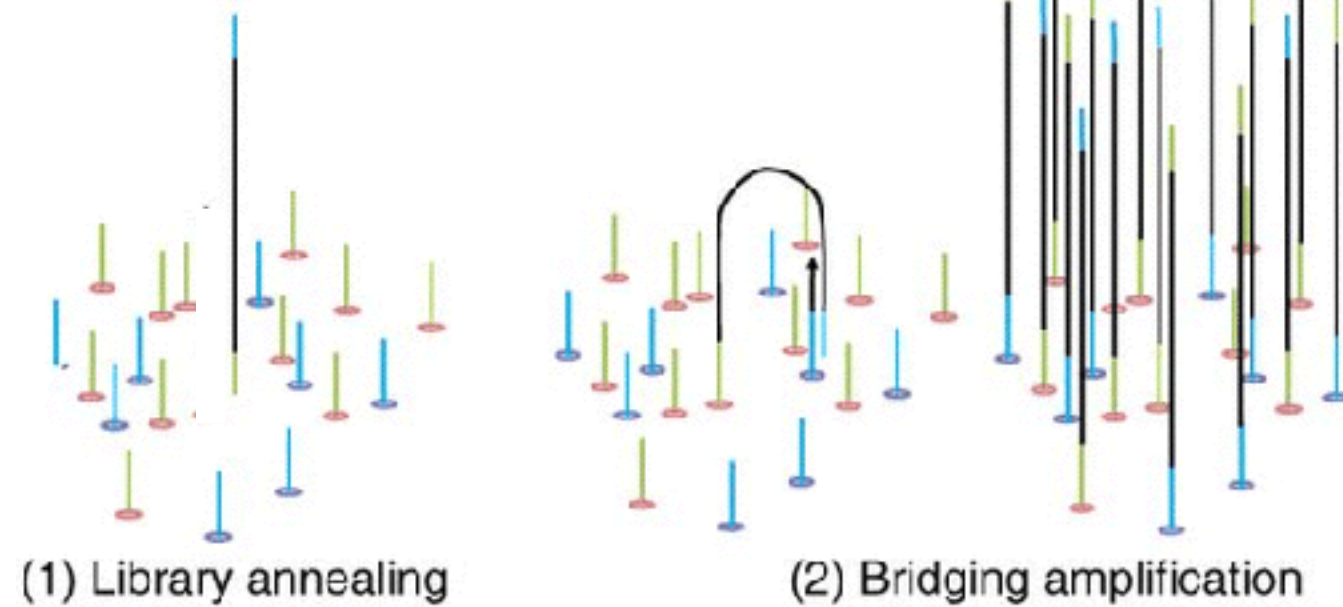


C Preparation for paired-end sequencing

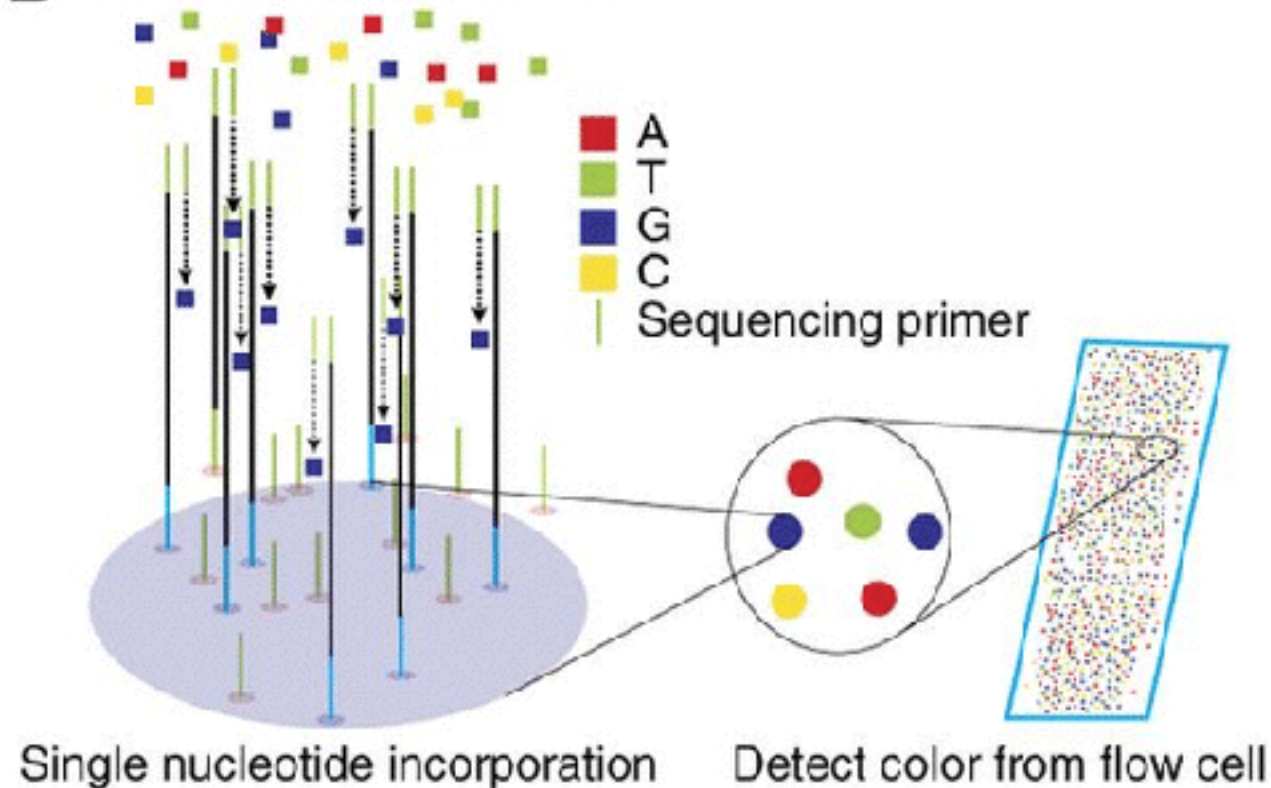


Pair-end Sequencing

A Cluster generation



B Sequencing by synthesis

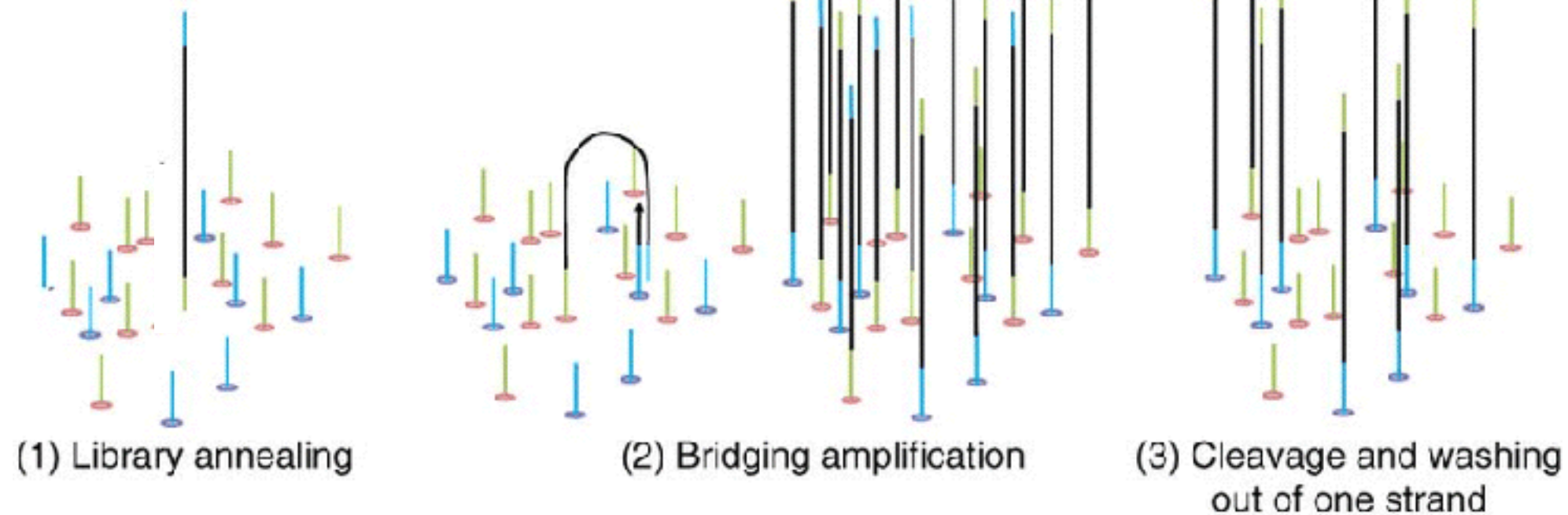


C Preparation for paired-end sequencing

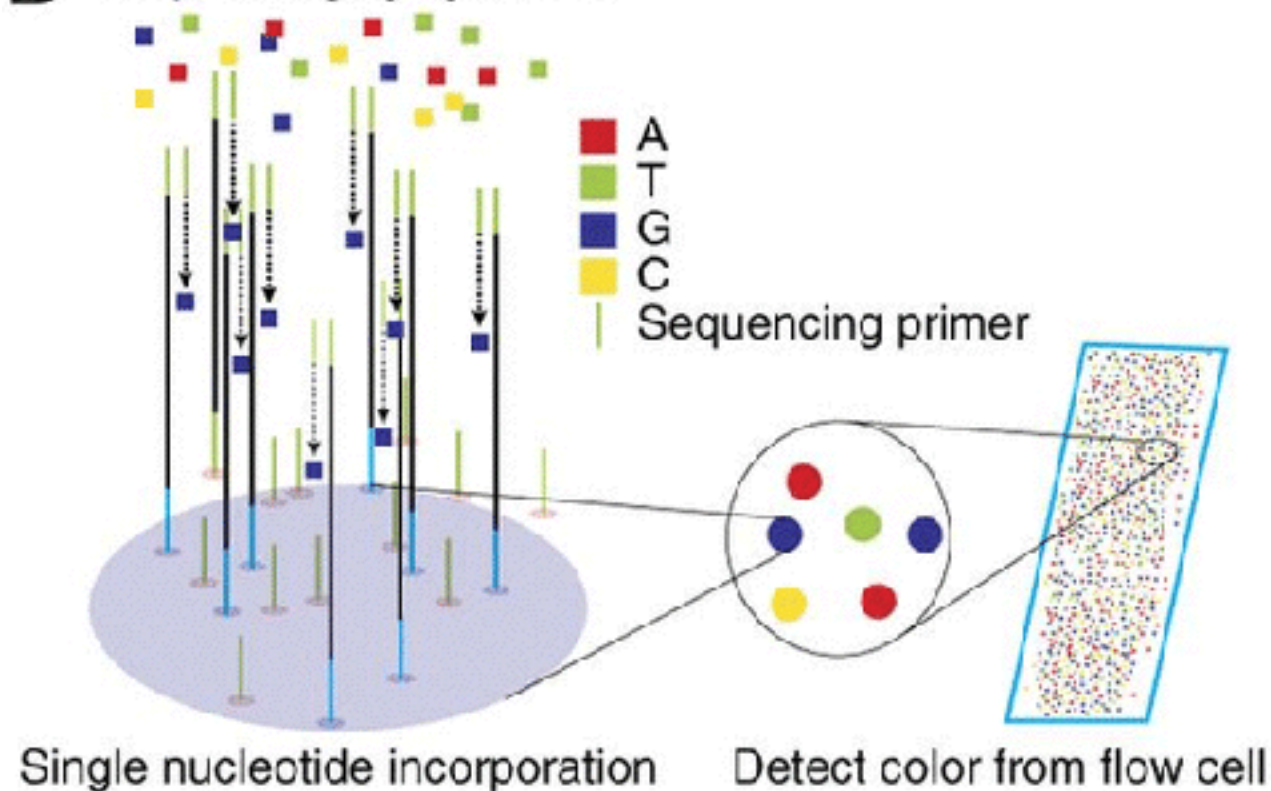


Pair-end Sequencing

A Cluster generation



B Sequencing by synthesis

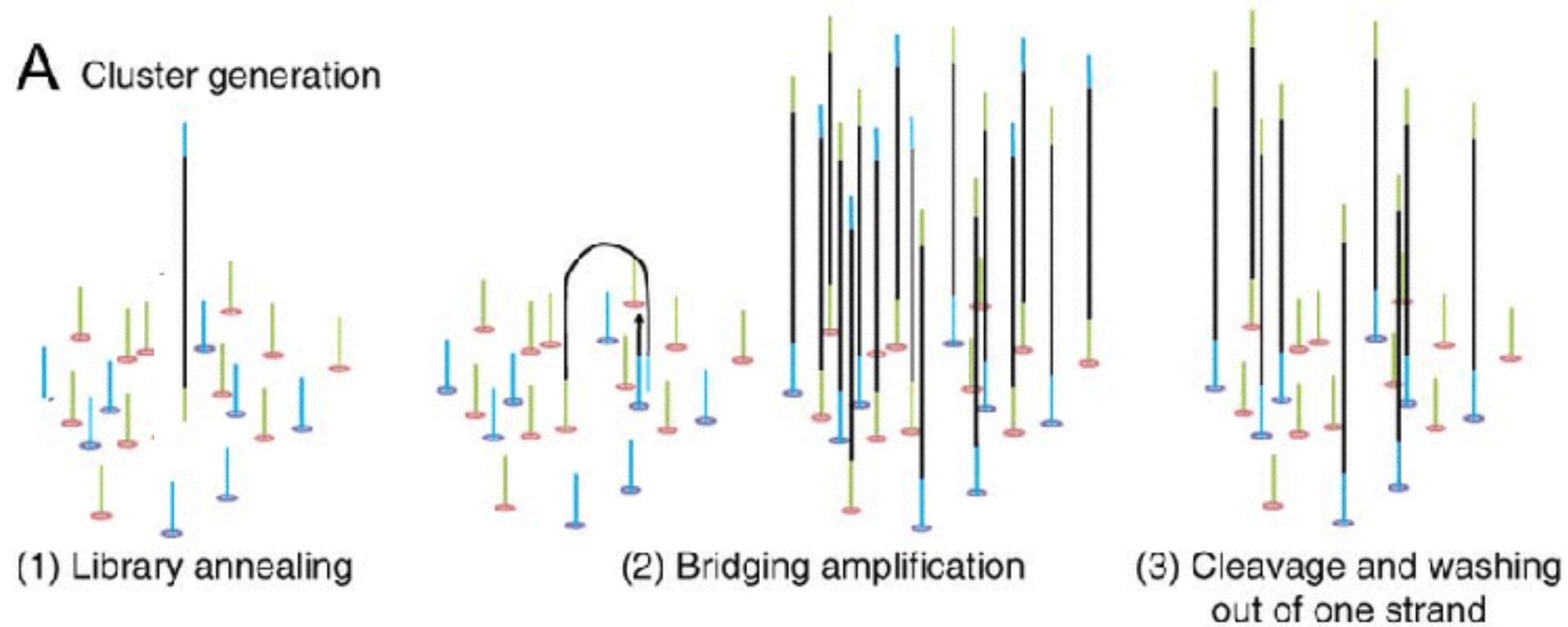


C Preparation for paired-end sequencing

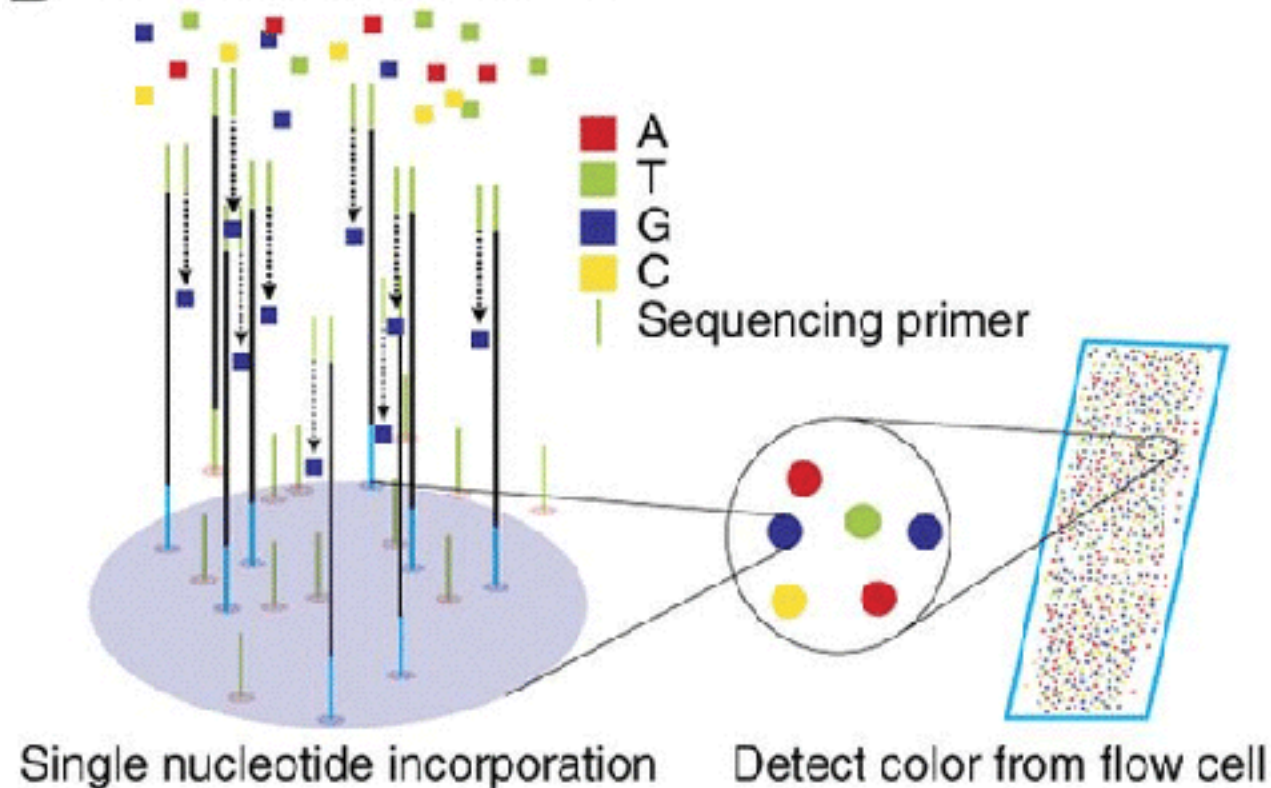


Pair-end Sequencing

A Cluster generation



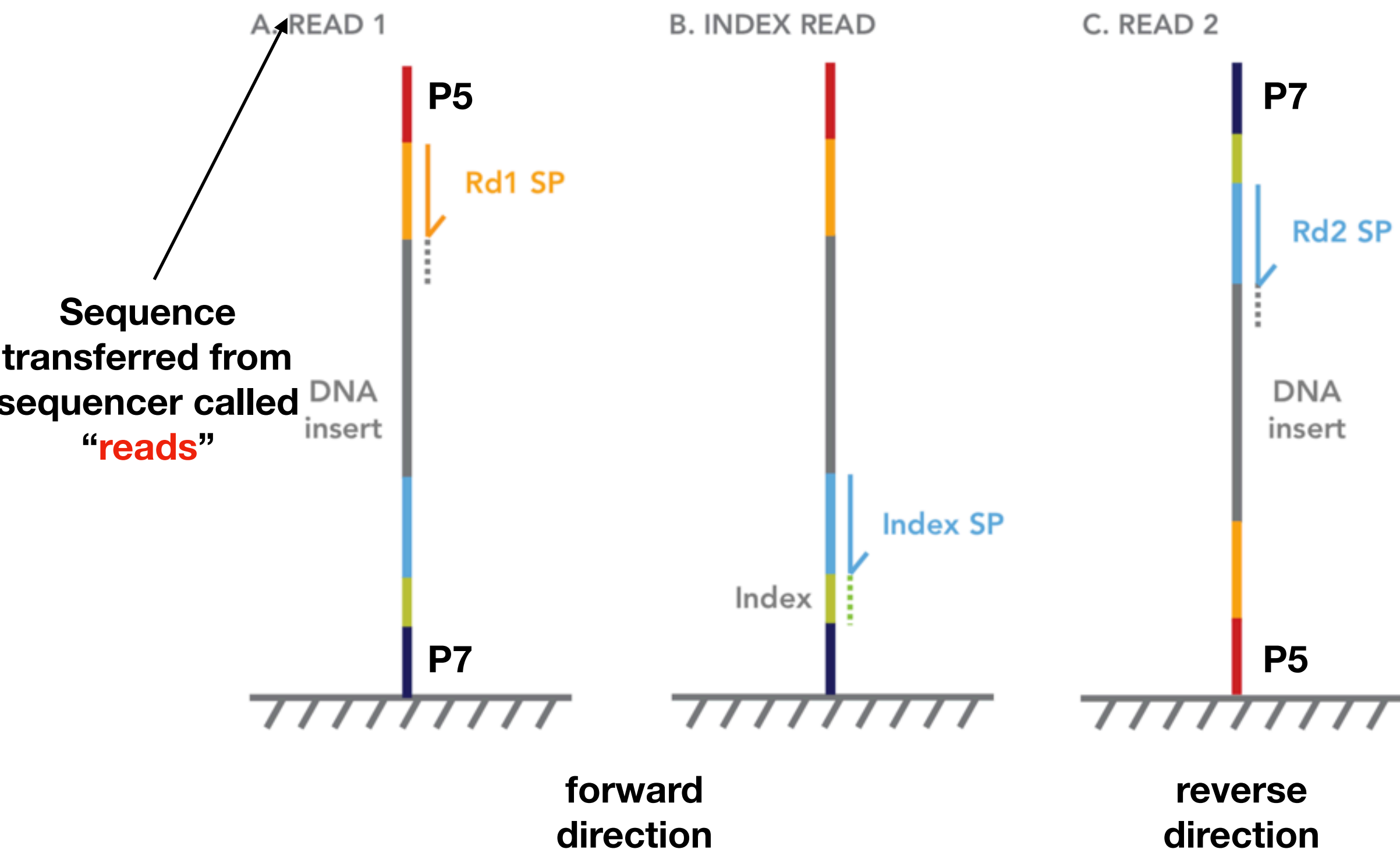
B Sequencing by synthesis



C Preparation for paired-end sequencing



Pair-end Sequencing



What comes from the sequencer

	Ai1_R1fq.gz
	Ai1_R2fq.gz
	Ai1.base.pdf
	Ai1.base.png
	Ai1.quality.pdf
	Ai1.quality.png

Assigned (more professionally called “demultiplexed”) reads according to index close to P7

Reads

**Bases
composition**

**Quality of
Bases**

File suffix is “fq” or “fastq”

↓

	AI1_R1fq
	AI1_R2fq

Expanded file

How reads in xxx_R1.fq looks like

header

index sequence

sequence

a “plus”.

quality (phred 33)

@E00487:317:HJLJHCCXY:5:1101:10734:2012(1):N:0:CGTCATAA
ACTACACTTTCTTTAGAAAATGCATTTTCTAGTTTATATTGATATCATTAAACATTTTGCCTAAACTAACTGGACTAACC
CTAATCAAATCATAAAATCTACAAGTTACCTAGTTTTTTGACCCGACATTAAAGGTGAGTCAA
+
FAJJJJJJJJJJJJJJFF-<FFJJJ<JJJJJJJJJJJJAFJJJJJJJJJJJJJJJJJJJJJJJJFJJJF<JFJ<FJFJJJFJFJJJ7
JFA-FAJFJJJJJJJJFFJJFJ<FFJJJJJJAJJFJJFFJJJJAAJ<FJ<JJJJJJJJJJJJJJJJFFJJFF

How reads in xxx_R2.fq looks like

header

index sequence

sequence

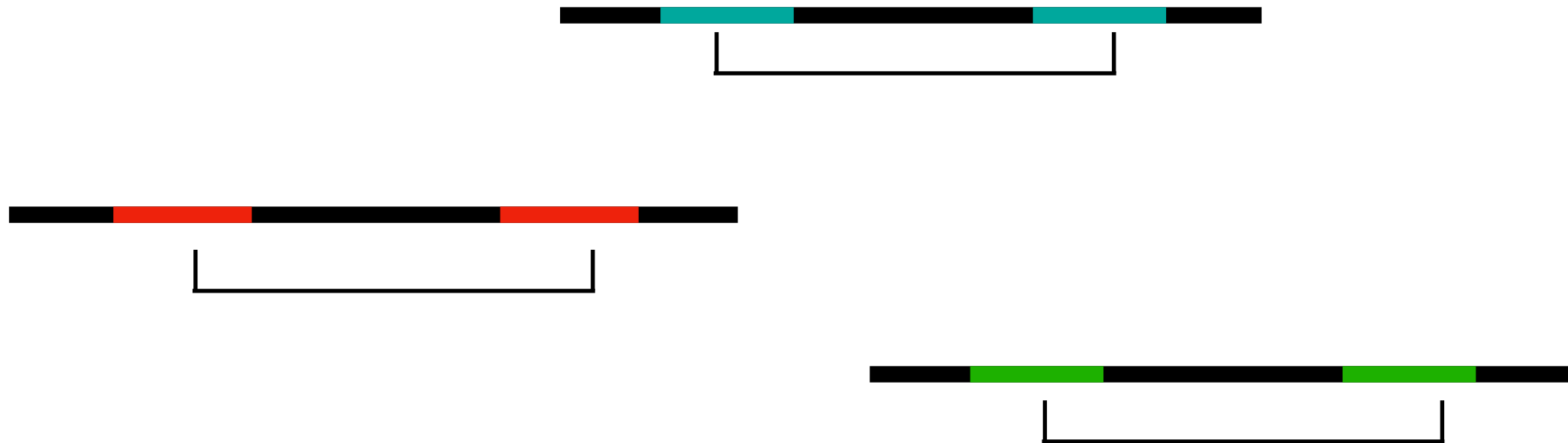
a “plus”

quality (phred 33)

```
@E00487:317:HJLJHCCXY:5:1101:20943:2012 2:N:0:CGTCATAA
AAGTTGACGGTGCTAGCAAGCTAAAGCTACCCGGAGATATGGAGGCAACTGAAGCCAACATTCTTGCGGAGCTGAAAAAA
GTACAGGATAATTTACAAGGTGAAATAGCGGACACGATTGGACTTGTGCAAACAGAAATGTCG
+
FFFFJJFFJJJJJJJA-AJJF-AFA<FJFJ<<AFJJ7-<F-FJJFJJ77--<F<FJ<-AF-<-<<-JJ77FJJ<J--AA---
-<7AFJJJJJ-7--7-FJ--<AFAJJJF-AF-77FAJ--7-7---7AF---7-<-AF7A7-7-
```

The diagram illustrates the structure of a FASTQ read. It shows a single read entry with five lines. The first line is the header, starting with '@'. The second line is the index sequence, which is circled in red. The third line is the sequence, followed by a plus sign on the fourth line. The fifth line is the quality score in phred 33 format. Arrows point from the labels to the corresponding parts of the read: 'header' to the first line, 'index sequence' to the second line, 'sequence' to the third line, 'a “plus”' to the plus sign on the fourth line, and 'quality (phred 33)' to the fifth line.

What's paired reads



**Reads originated from the same molecular before bridge amplification
called paired reads**

Paired reads in xxx_R1.fq and xxx_R2.fq

xxx_R1.fq

[illegible]

xxx_R2.fq

@E00487:317:HJLJHCCXY:5:1101:20943:2012 2:N:0:CGTCATAA
AAGTTGACGGTGCTAGCAAGCTAAAGCTACCCGGAGATAATGGAGGCAACTGAAGCCAACATTCTTGCGGAGCTGAAAAAA
GTACAGGATAATTTACAAGGTGAAATAGCGGACACGATTGGACTTGTGCAAACAGAAATGTCG
+
FFFFJJFFJJJJJJA-AJJF-AFA<FJFJ<<AFJJ7-<F-FJJFJJ77--<F<FJ<-AF-<-<-JJ77FJJ<J--AA---
-<7AFJJJJJ-7--7-FJ--<AFJJJJF-AF-77FAJ--7-7---7AF---7-<-AF7A7-7-

Paired reads have unique id in their corresponding file and exactly the same name

Compare reads in xxx_R1.fq and xxx_R2.fq

xxx_R1.fq

```
1 @E00492:247:HFMH3CCXY:8:1101:18436:2909 1:N:0:CTGCAATG
2 CTCAAAAACAACAATAATCTATTATCACAACCAAAGCAACAGCATACATATAATATCAATGCTTCTTTATCTCGATGTGAATAATAAGAAGTCAATGGGAACCTTACCTTT
3 +
4 FJJJJAJJJJJFJJ<-FJJJFJJJJJJFJJJJFJJFJJJJJJJJAJJJJJJJJJJJFJJJJ<JJJJJJFFAFJJF77-A<AJJJJJAJJJJJFFJJJJ7<F-7FJJ-AJF-F
5 @E00492:247:HFMH3CCXY:8:1101:6725:3401 1:N:0:CTGCAATG
6 AATTCAGCGACTTGGTTCACACAATCCTCCATCCCTCCTTCTTTACCACGATGTTTGATGATGCTGCAGGTACAGCTGCTGAAGAGACACGAGCTGTGACCAGTAACTA
7 +
8 7JJJJJJJJ-FJJJJJJFJJJJAJFJJJA-FFJJFJJJJAAJJF7AJJJJJJJF7A7JJJJFJJJJFJJJJAFJJAJFFFJJJJ7JJFFJJFJJJJJJF<-----77---7--77
9 @E00492:247:HFMH3CCXY:8:1101:8745:3841 1:N:0:CTGCAATG
10 AAATGTAAGGTCAAACCTTGTCTATGAAGCCGGTGTGCCATTTGCGTAGCTTCCTGTTTTCGGTGGAGAGTTTCTCAGCAGCAGACTGCAGCTTGGCATTGTCA
11 +
12 JJJJJJJJJJJJJJJJJJJFJJJJFFJJJJJJFJJJJJJJJJJJJFAFJJJJJJ<FJJJJFJJJJJFA<FJ-<AAJFJJJJF-7A-AJJ-7JA7FJJFJJJJJJJJJJJJ7JJ7777--7-<A<7A
```

xxx_R2.fq

```
1 @E00492:247:HFMH3CCXY:8:1101:18436:2909 2:N:0:CTGCAATG
2 ATAATGGTTTTTGGACAAGTTTAGCTTAAAGGAGGAAAATTGCATGTGTTGTGATCTTGTCACTCTCTGTGTCTCCTAGACTCCTGCAGAATTGCAAGGAGGACAGGTTAGA
3 +
4 AJJ--7F<-<F---<F--FFJJ7FJ-77AJ-AF-77-7F<-<<---F-7-AF--7---<-----7AAA7-7-----7AJ-A-AA-A-AJF7A7--AAF7J<A7FJAJ--7A-F
5 @E00492:247:HFMH3CCXY:8:1101:6725:3401 2:N:0:CTGCAATG
6 GGTATCTTTAAAACACATTTAACATGCGATACATAACTTATAATTGGTCATGTTTGTGCATACAAATTGAATTAGACAGTGAATCAAAGAAAATGTGCTTACACTGACAAGT
7 +
8 JAA7A--<FJJ--<JFJJ7FF<JFJJ<-F-<JJ-FJJF77F<FJJJJFJJJA-FJJJJFJ<<JJJJJJFFF<F-FFA-A--7<-F-AFAJJJA<FFJJJFFFF-JJJFJF--A7--77-
9 @E00492:247:HFMH3CCXY:8:1101:8745:3841 2:N:0:CTGCAATG
10 AATGGTTGCCAAGCAACACAGAGACGAAGTGGCAAGTATGATTGAAGATTGGAGTGATATGAGTGATATTAACATTTCAGCTGGGGGTGATTATTAACATTTAGCAGGATAT
11 +
12 FAF<-7FA<7-<7<<7FFJJ7F<<--A--<AFF-7F-A<FF-<A<-FF7FF7AJF77<JF-<AF-J-F-AFAFJJJAF-7<AAAAAJAFJ<FFFJJJJJJF<J<<A<<<-A-7
```

Paired reads are placed at the same line of its corresponding file

Let's start analysis now!!

What's the goal of analysis?

The sequence of loci used to design the baits, we call it “**reference**”

locus 1



locus 2



locus 3



Intact sequences of corresponding **locus** of each sample

locus 1



locus 2



locus 3



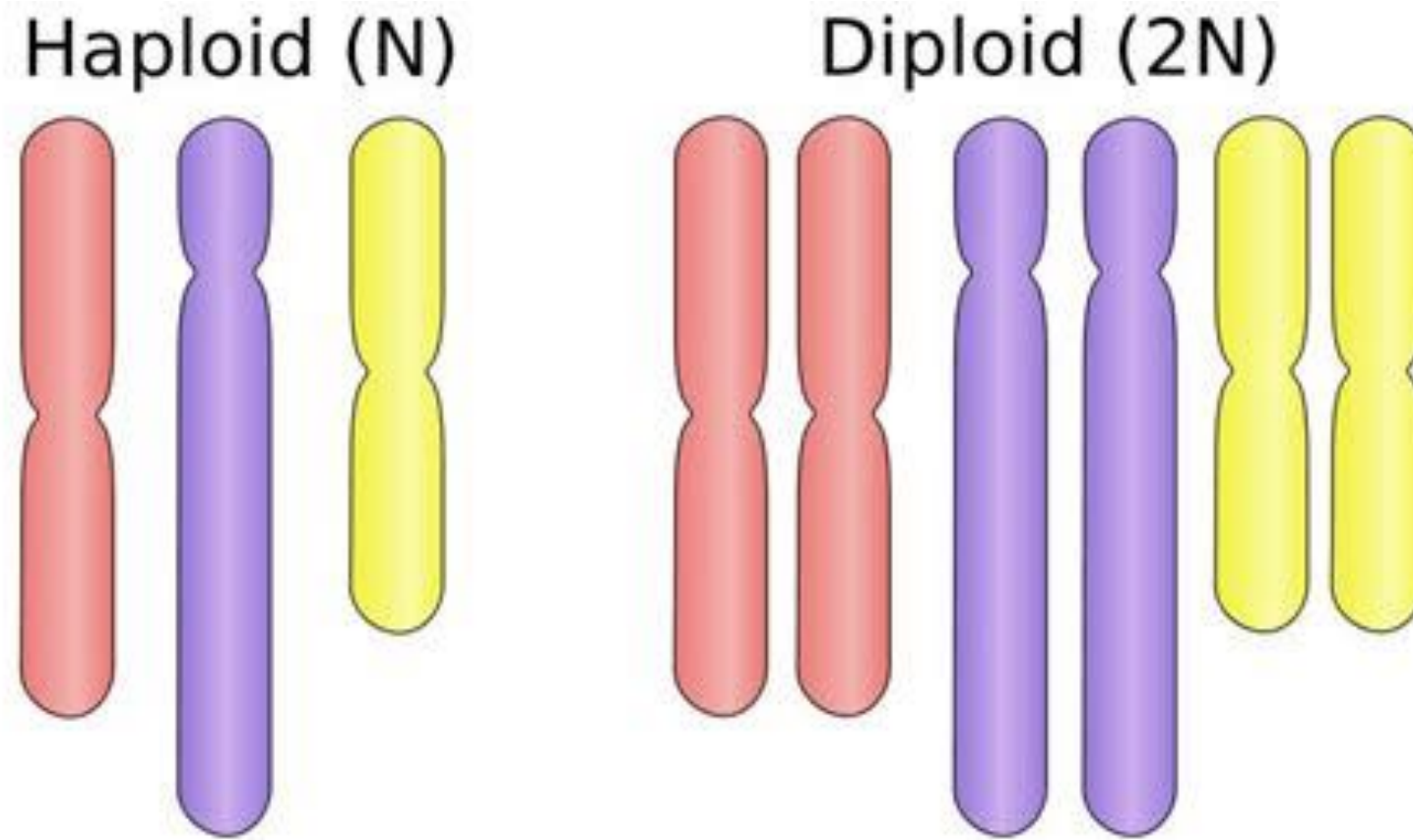
A **locus** is a fixed position on chromosome. In our lab, a locus is always an exon flanked by intron or UTR region

What's the goal of analysis?

1	@E00492:247:HFMH3CCXY:8:1101:18436:2909 1:N:0:CTGCAATG	
2	CTCAAAAACAACAATAATCTATTATCACAACCAAAGCAACAGCATACATATAATATCAATGCTTCTTTATCTCGATGTGAATAATAAGAAGTCAATGGGAACCTTACCTTT	
3	+	
4	FJJJJAJJJJJFJJ<-FJJJFJJJJJJFJJJJFJJFJJJJJJJJAJJJJJJJJJJJFJJFJJ<JJJJJJFFAFJJF77-A<AJJJJJAJJJJJFFJJJJ7<F-7FJJJ-AJJ-F	
5	@E00492:247:HFMH3CCXY:8:1101:6725:3401 1:N:0:CTGCAATG	
6	AATTCAGCGACTTGGTTCACACAATCCTCCATCCCTCCTTCCTTTACCACGATGTTTGATGATGCTGCAGGTACAGCTGCTGAAGAGACACGAGCTGTGACCAGTAACTA	
7	+	
8	7JJJAJJJ-FJJJJJFJJJJAJFJJJA-FFJJFJJJJAAJJF7AJJJJJJJF7A7JJJJFJJJJFJJJAFJJAJFFFJJJJ7JFFJJFJJJJJJF<-----77---7--77	
9	@E00492:247:HFMH3CCXY:8:1101:8745:3841 1:N:0:CTGCAATG	
10	AAATGTAAGGTCAAACCTTGTCTATGAAGCCGGTGTGCCATTTGCGTAGCTTCCTGTTTTCGGTGGAGAGTTTCTCAGCAGCAGACTGCAGCTTGGCATTGTCA	
11	+	
12	JJJJJJJJJJJJJJJJJFJJJJFFJJJJJFJJJJJJJJJJJJFAFJJJJJJ<FJJJJFJJJJJFA<FJ-<AAJFJJJJF-7A-AJJ-7JA7FJJFJJ7JFJJJJJJ7J7777--7-<A<7A	

Short raw data

There's only one sequence for each sample



locus 1

sample1 

sample2 

sample3 

3 steps to recover qualified assemblies from raw data

Data preparation

Assembling

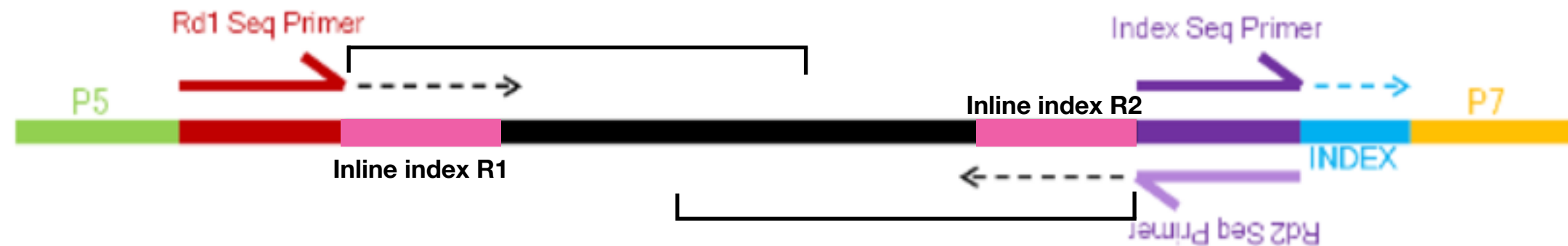
Further processing

Data preparation

Demultiplex reads according to inline index

Trim low quality bases and adaptor

Demultiplex reads according to inline index



Reads we got still includes inline index

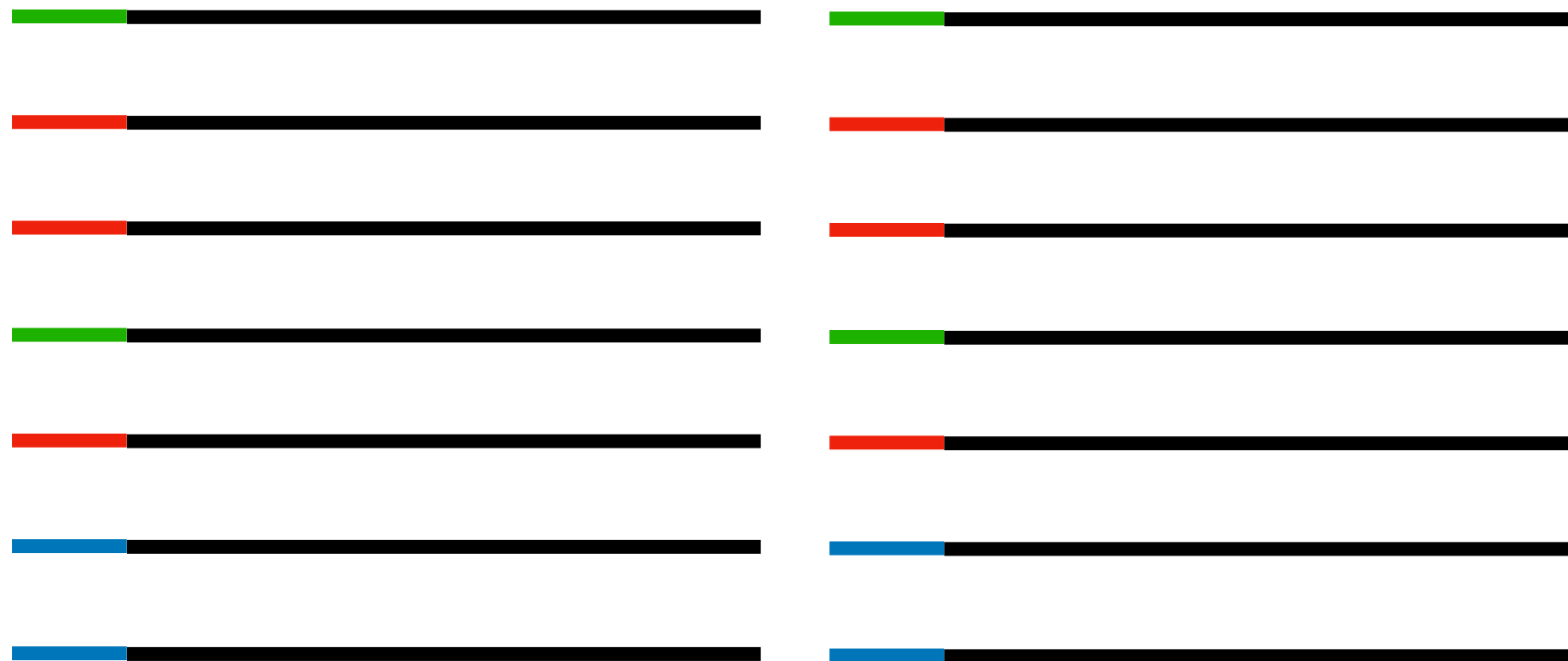
We need to **demultiplex** reads according to them, then cut them out

←
assign reads to its sample

How to demultiplex

Reads before demultiplexing

Paired Reads



Table

Index “**green**”: sample 1

Index “**red**”: sample 2

Index “**blue**”: sample 3

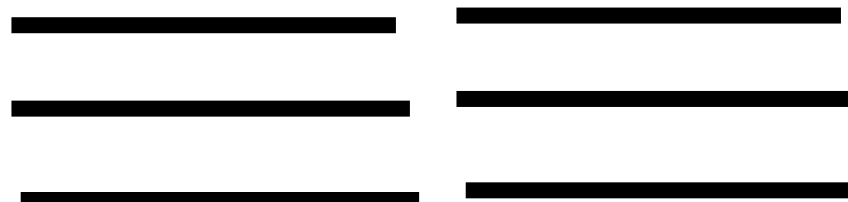


Reads after demultiplexing

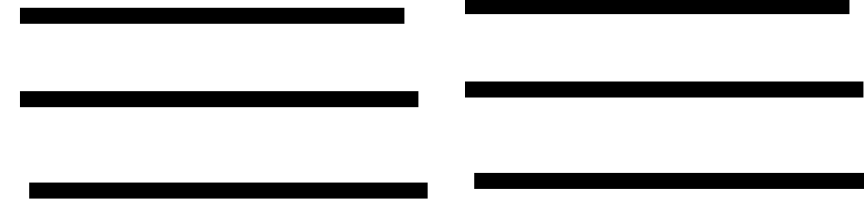
sample 1



sample 2



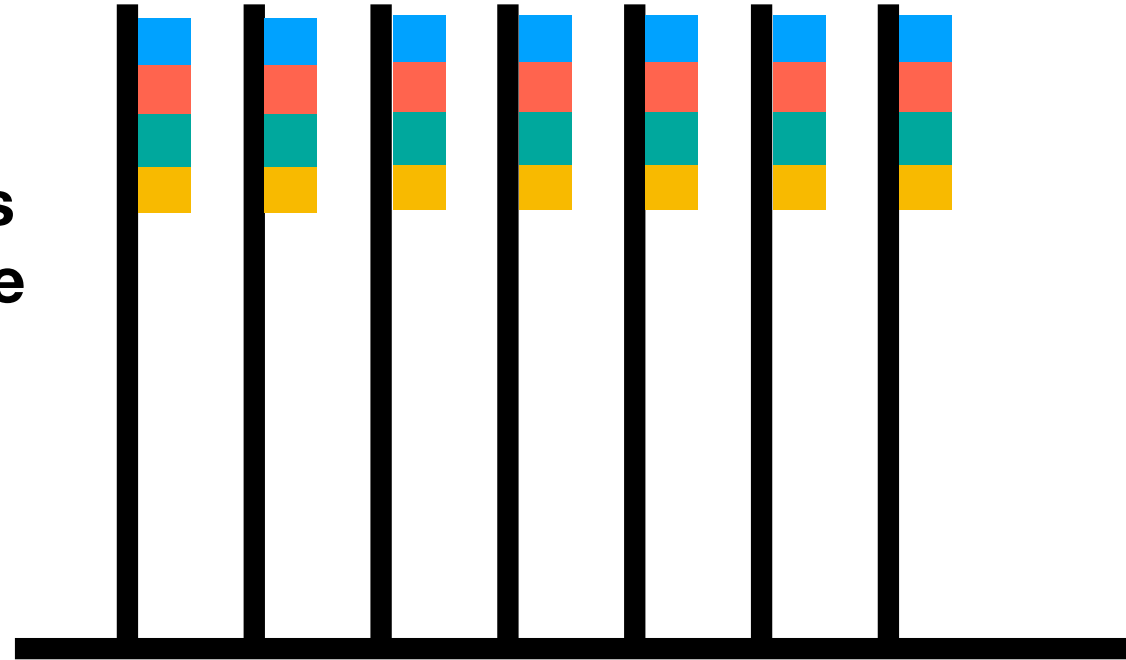
sample 3



Trim low quality bases and adaptor

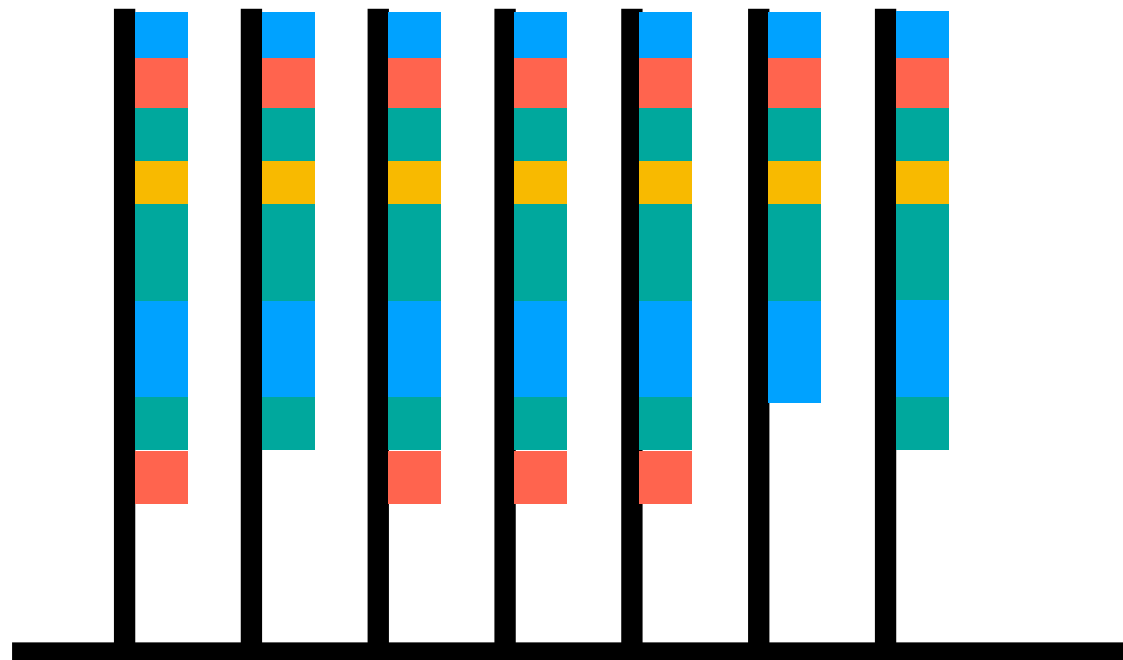
Why we need to trim low quality bases

A cluster of reads
amplified from the
same molecular



Last base is 

Several rounds of extension later

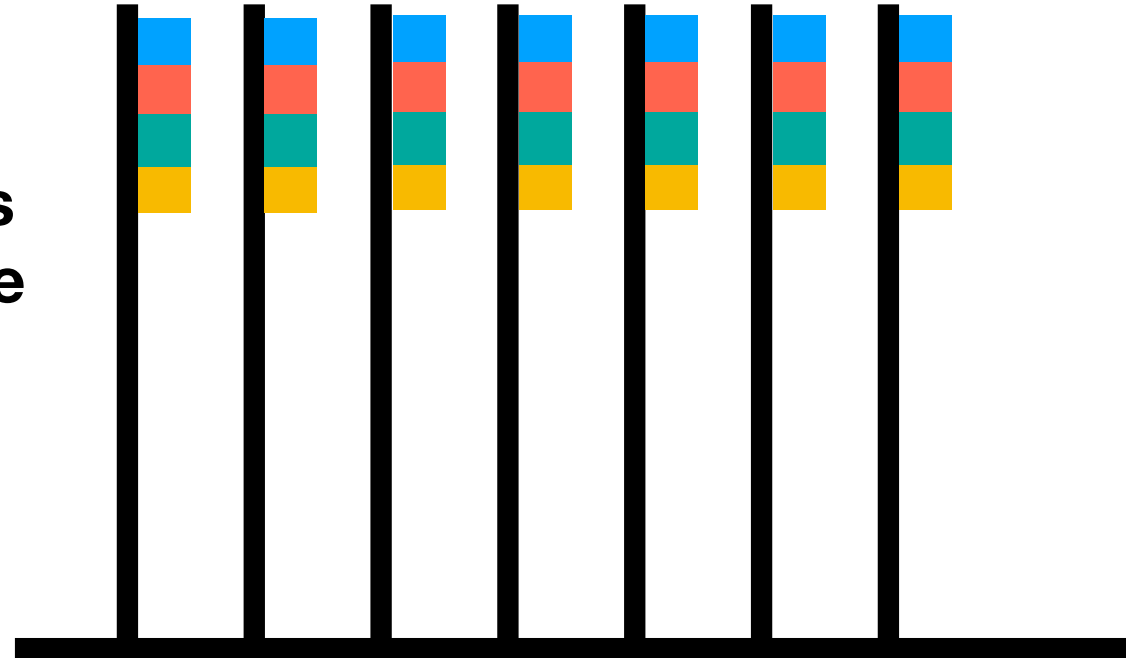


Last base is  or  ?

Trim low quality bases and adaptor

Why we need to trim low quality bases

A cluster of reads
amplified from the
same molecular



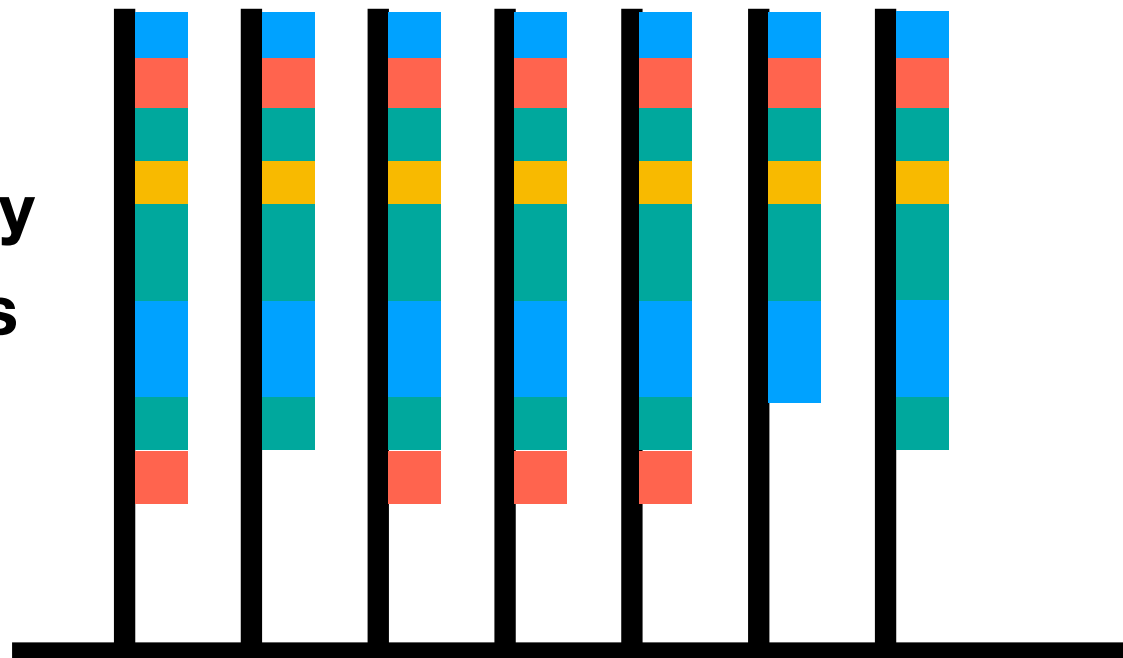
Last base is 

High quality bases

Several rounds of extension later



longer the time of
sequencing quality
of bases becomes
lower



Last base is  or  ?

Low quality bases

How to trim low quality bases

$E(\text{quality}) \geq 15$

ACGGCGTAGGCTGATGATCGGGGTACGTCCGATCGTAGCTGTCA

$E(\text{quality})=30$ GOOD

ACGGCGTAGGCTGATGATCGGGGTACGTCCGATCGTAGCTGTCA

$E(\text{quality})=29$ GOOD

ACGGCGTAGGCTGATGATCGGGGTACGTCCGATCGTAGCTGTCA

$E(\text{quality})=25$ GOOD

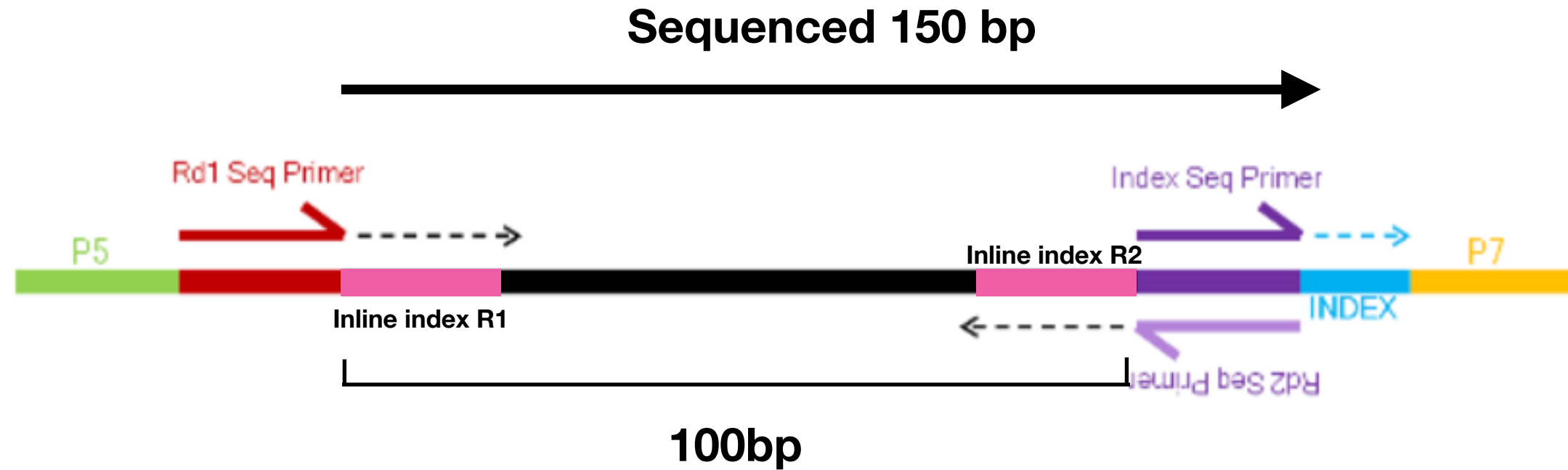
...

ACGGCGTAGGCTGATGATCGGGGTACGTCCGATCGTAGCTGTCA

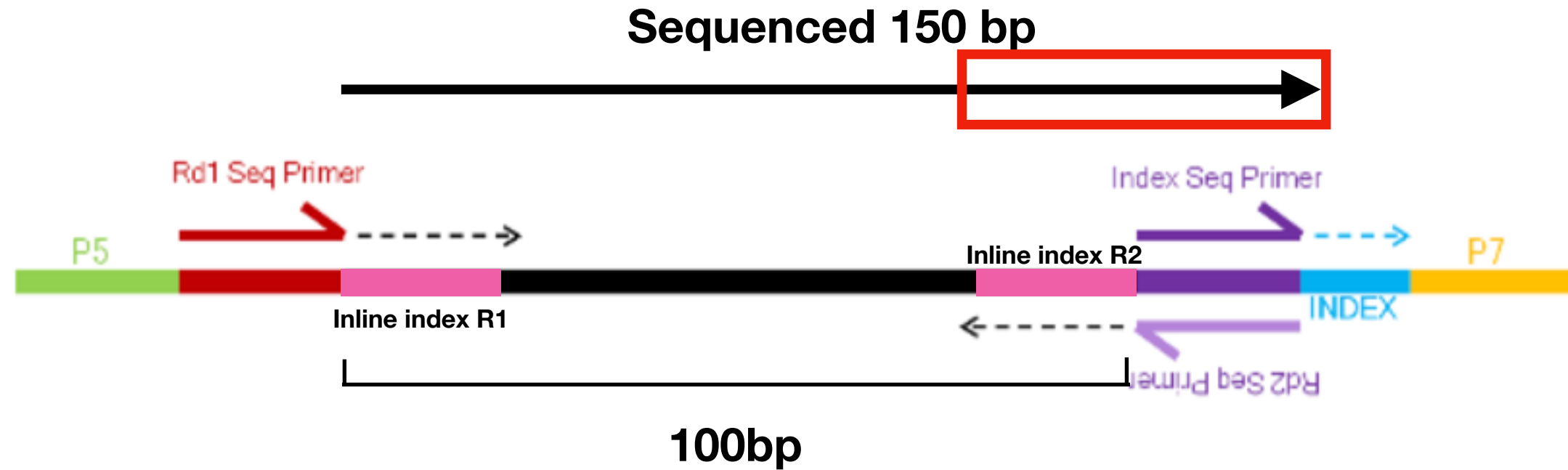
$E(\text{quality})=10$ BAD!!

ACGGCGTAGGCTGATGATCG

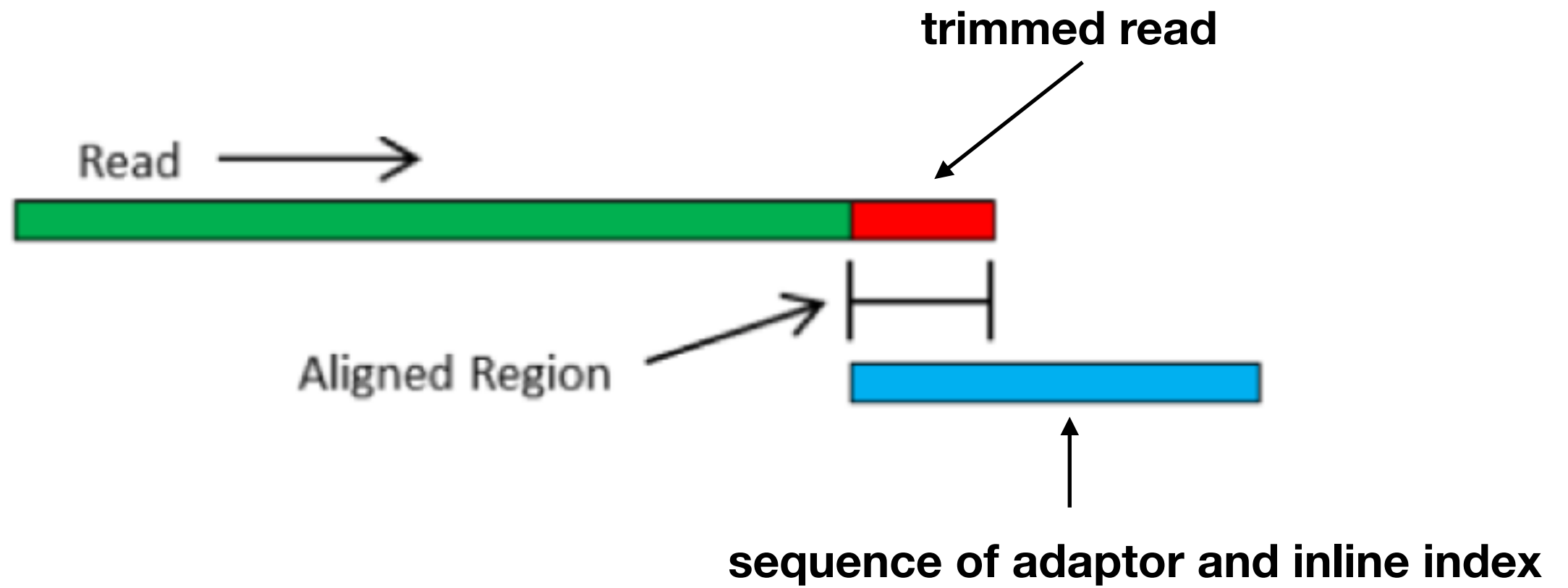
Why we need to trim adaptor?



Why we need to trim adaptor?



How to trim adaptor?



Reads have been cleaned. Let's start **assemble !**

What's assemble

Reads 1: CGGCGGATCTGATGGGATCTGATTCGGTT

Reads 2: TCTGATTCGGTTCGGATCTGGGCAT

Reads 3: ATCTGGGCATGGCGTTCGATGTCGCTAT

3 reads in a sample

What's assemble

Reads 1: CGGCGGATCTGATGGGAT**TCTGATTCGGTT**

Reads 2: **TCTGATTCGGTT**CGGATCTGGGCAT

Reads 3: ATCTGGGCATGGCGTTCGATGTCGCTAT

Resulting contig:

Contig1 CGGCGGATCTGATGGGAT**TCTGATTCGGTT**CGGATCTGGGCAT

“**Contig**” is the sequence assembled from the reads

What's assemble

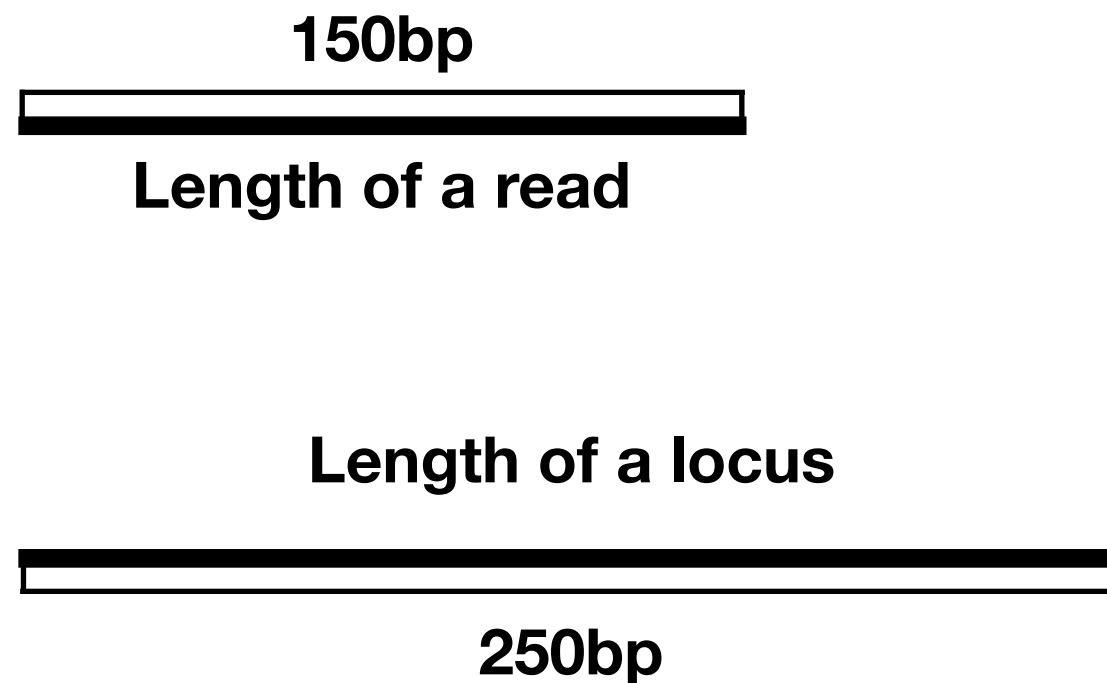
Contig1 CGGCGGATCTGATGGGATCTGATTCGGTTCGG**ATCTGGGCAT**

Reads 3: **ATCTGGGCAT**GGCGTTCGATGTCGCTAT

Resulting contig:

CGGCGGATCTGATGGGATCTGATTCGGTTCGG**ATCTGGGCAT**GGCGTTCGATGTCGCTAT

Why raw data need to be assembled before various analysis?



Reads are too short to reach the length of the locus

How raw reads magically become sequences of loci of each sample?

Remove PCR duplicates

Parse reads to loci

Assemble parsed reads

Further assemble

Get orthologue assemblies

Remove PCR duplicates

Reads 1: CGGCGGATCTGATGGGAT**TCTGATTCGGTT**

PCR duplicate of Reads 1: CGGCGGATCTGATGGGAT**TCTGATTCGGTT**

PCR duplicate of Reads 1: CGGCGGATCTGATGGGAT**TCTGATTCGGTT**

Reads 2: **TCTGATTCGGTT**CGGATCTGGGCAT

Resulting contig:

Contig1 CGGCGGATCTGATGGGAT**TCTGATTCGGTT**CGGATCTGGGCAT

Remove PCR duplicates

Reads 1: CGGCGGATCTGATGGGAT**TCTGATTCGGTT**

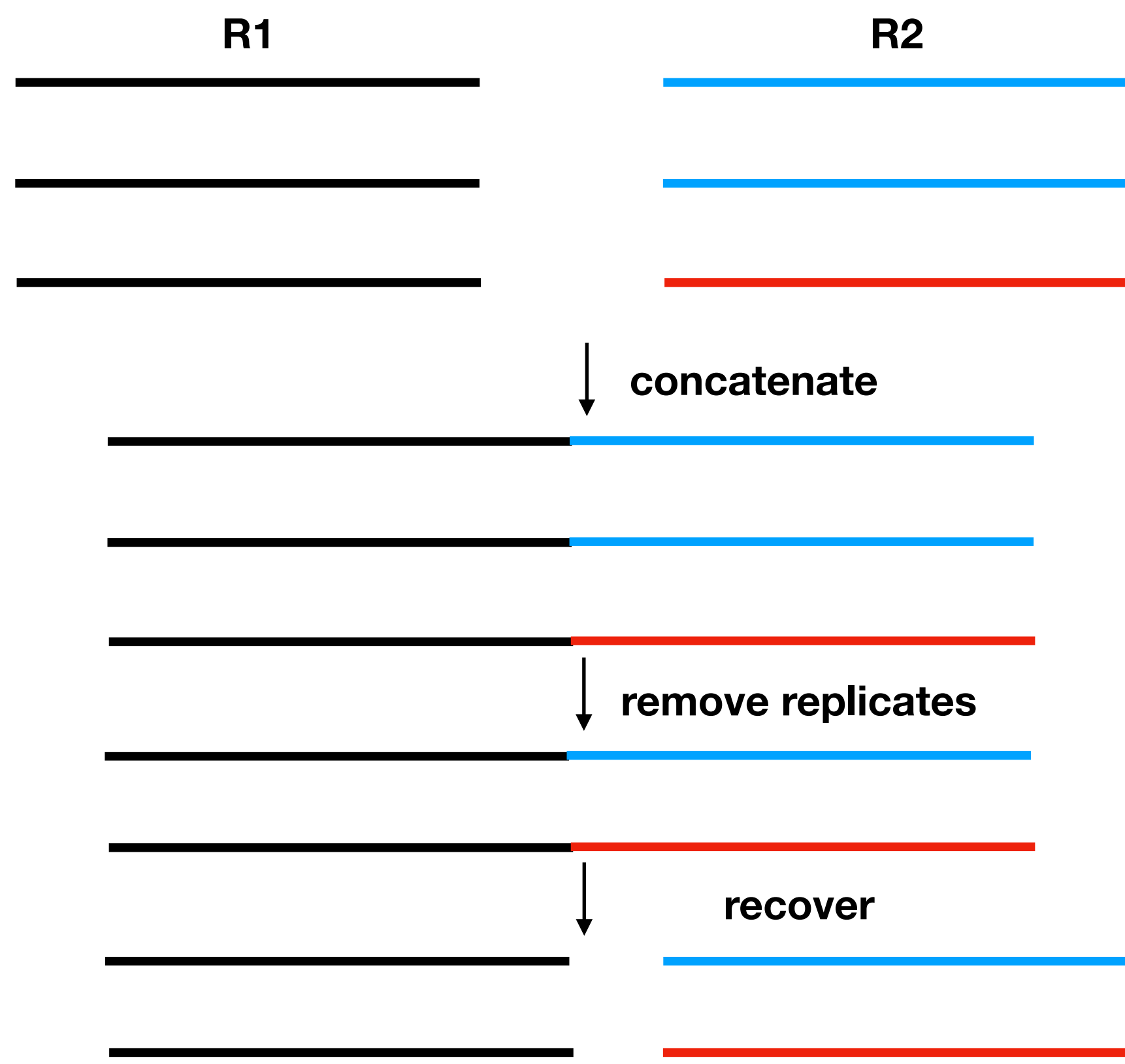
Reads 2: **TCTGATTCGGTT**CGGATCTGGGCAT

Resulting contig:

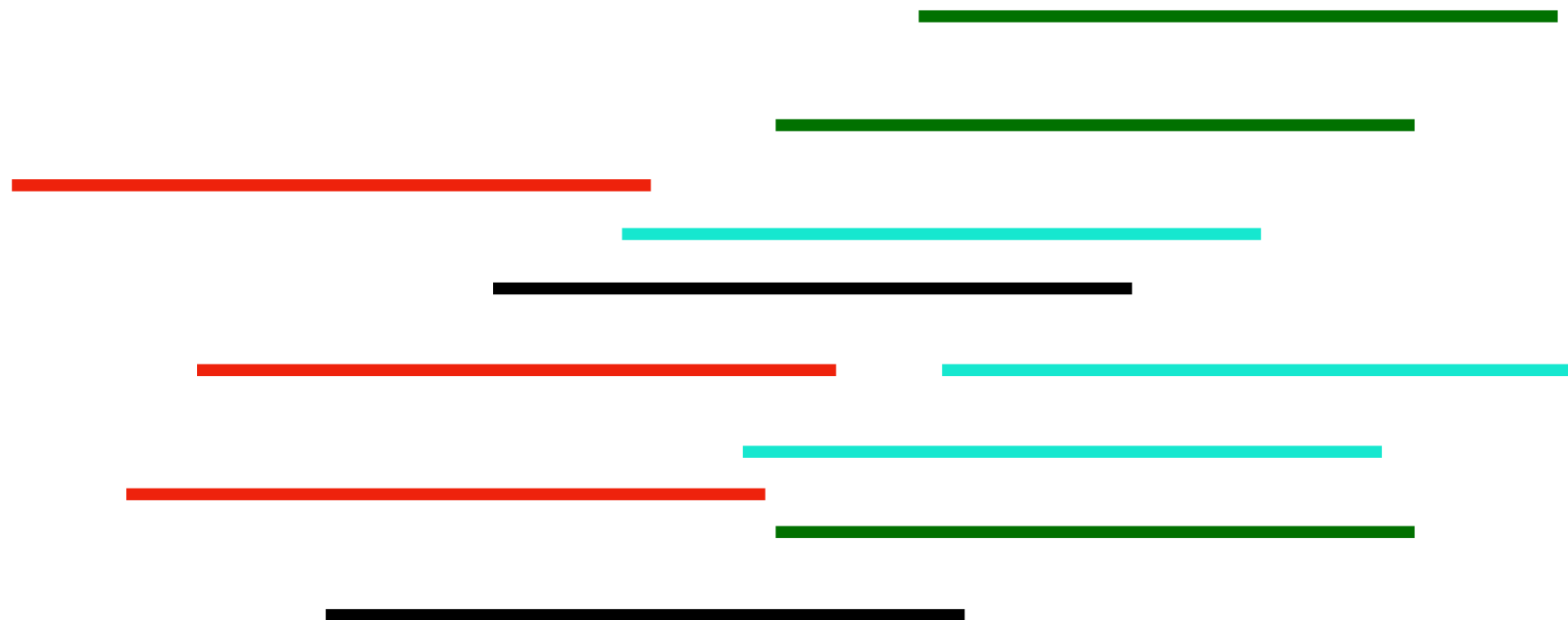
Contig1 CGGCGGATCTGATGGGAT**TCTGATTCGGTT**CGGATCTGGGCAT

PCR duplicates are redundant for assembly

Remove PCR duplicates



Parse reads to loci

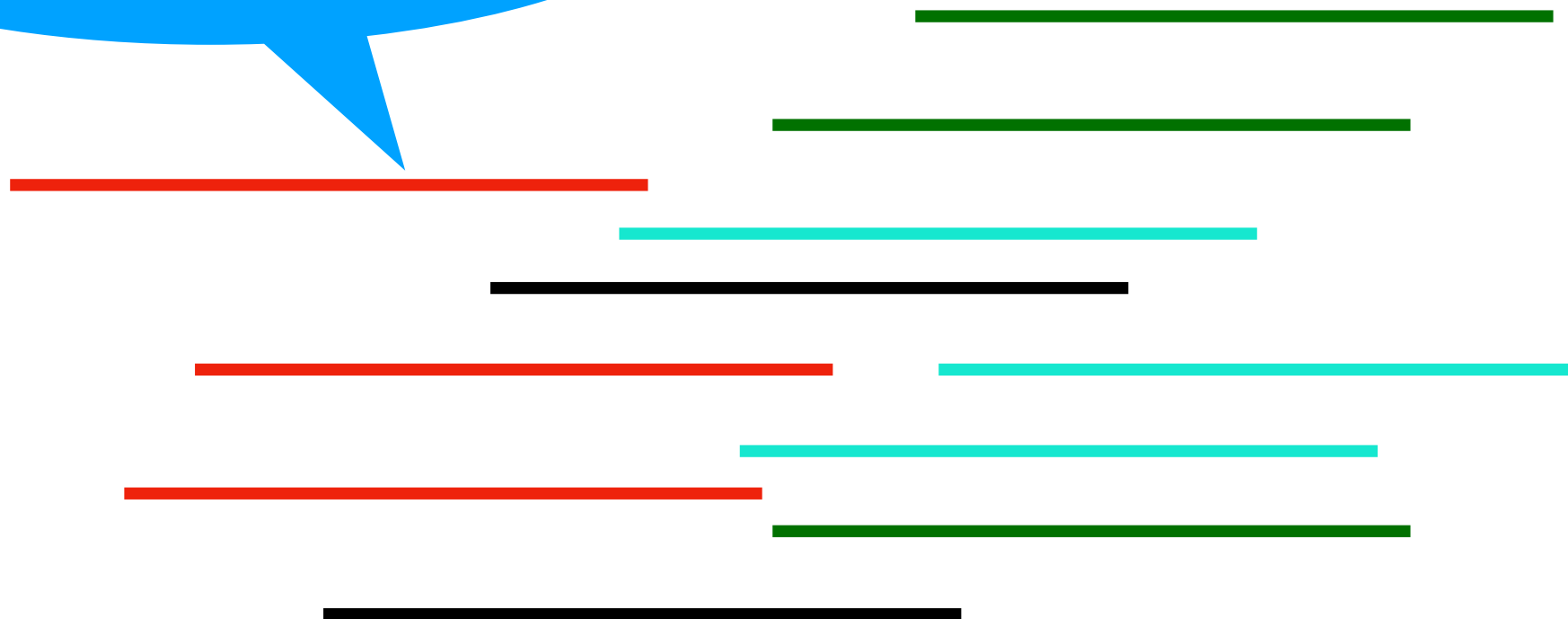


Mixed short reads from lots of loci

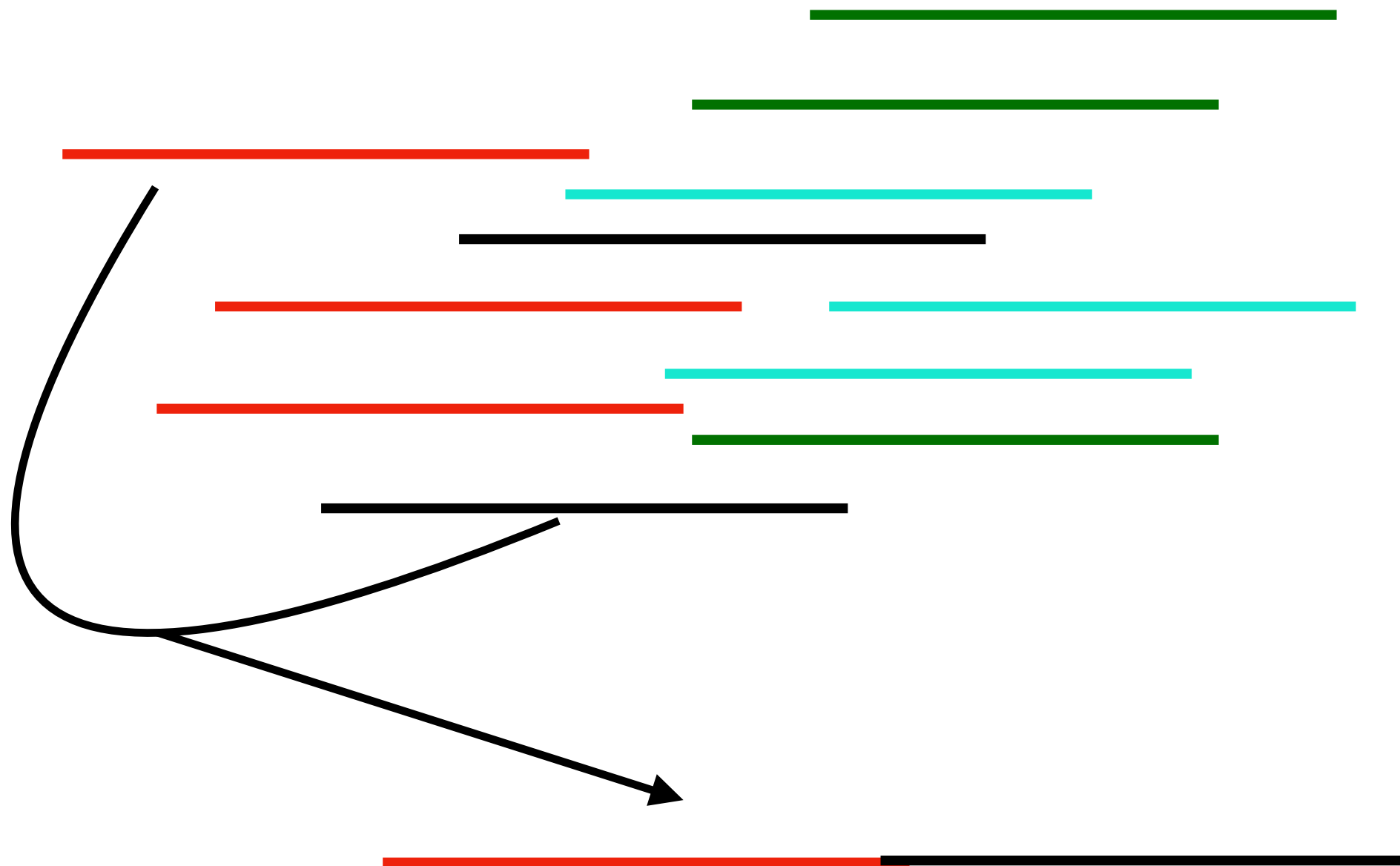
Reads of the same color indicate they come from the same loci

Parse reads to loci

I should assembled with which reads?



Parse reads to loci



If reads from different loci assembled together,
the resulting contig will be “**chimera**”

Parse reads to loci

locus 1



locus 2

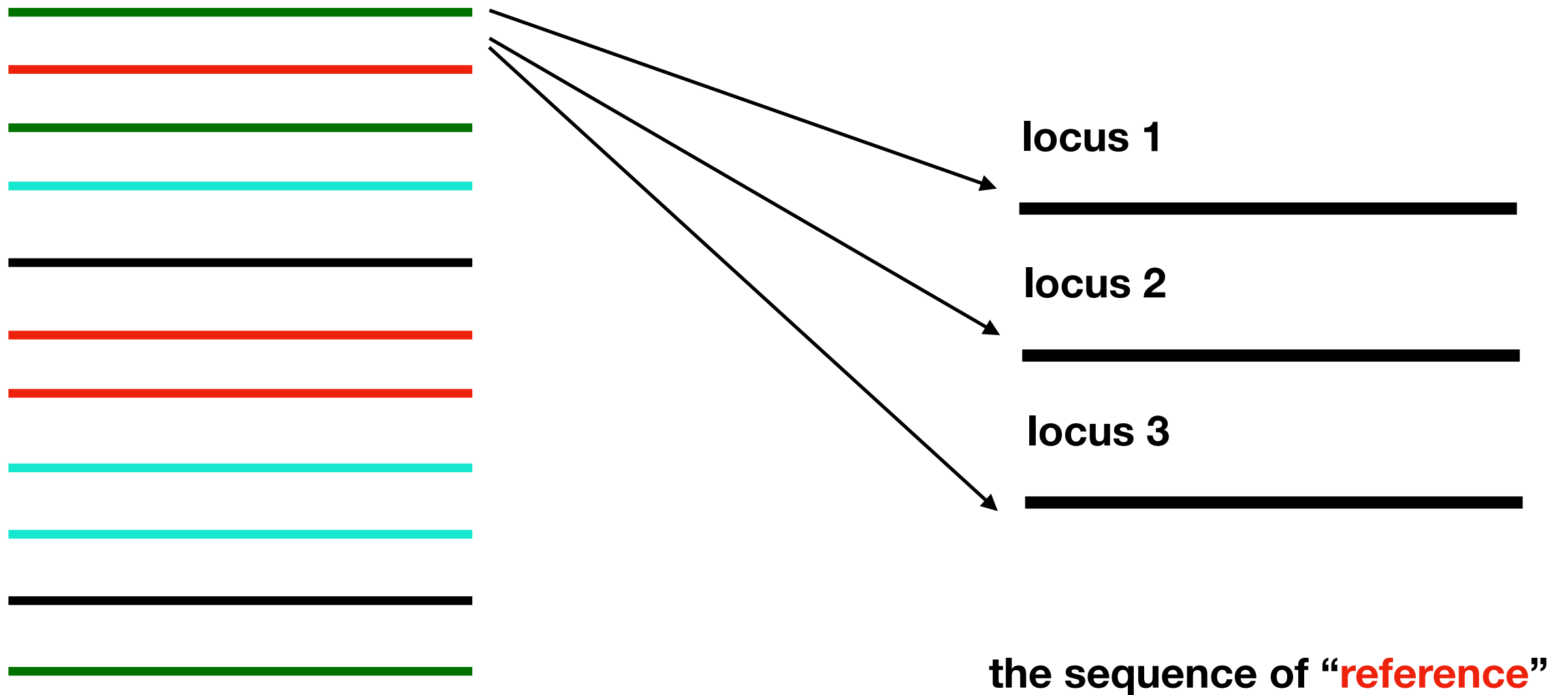


locus 3



Remember me? I'm the sequence of
“**reference**”, used to design the baits

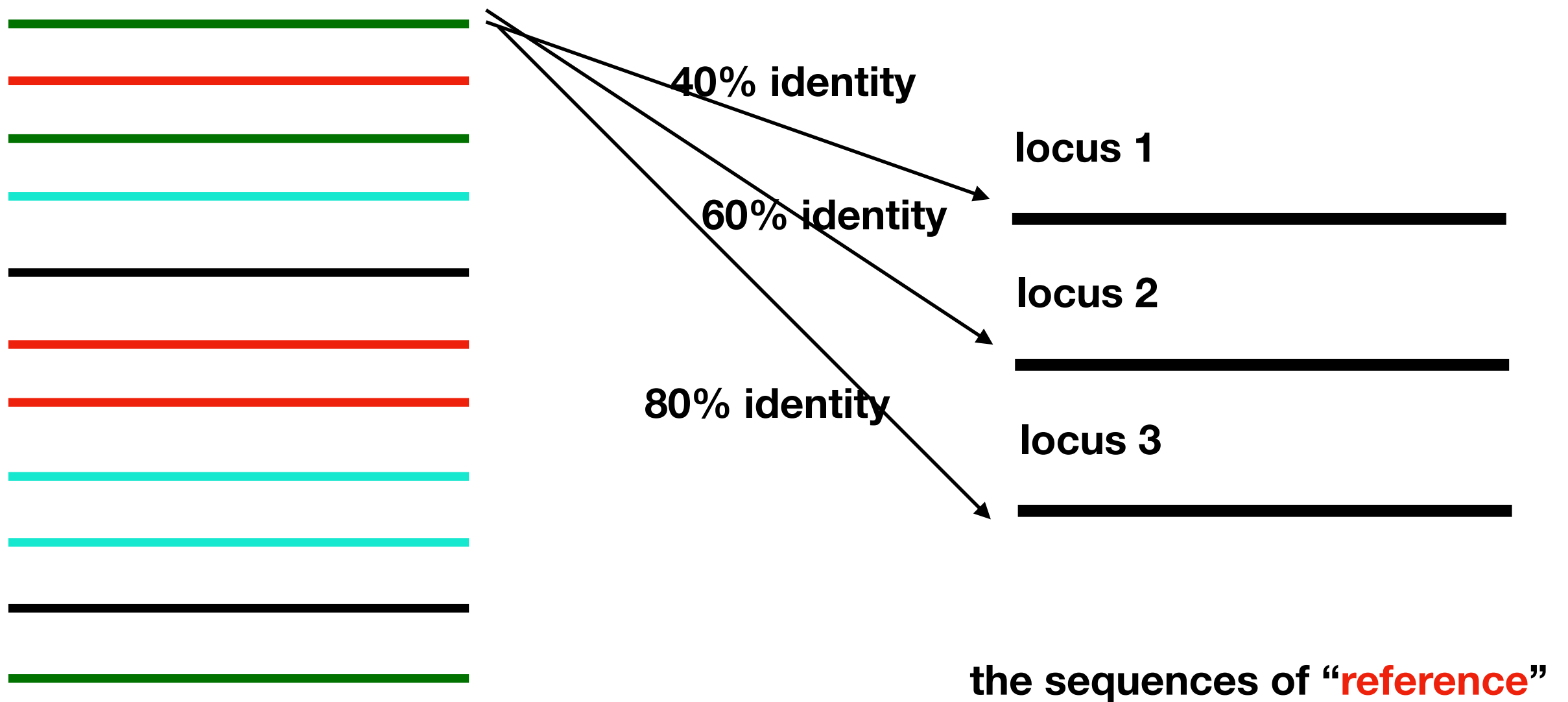
Parse reads to loci



Mixed short reads from several loci

Compare reads with each locus

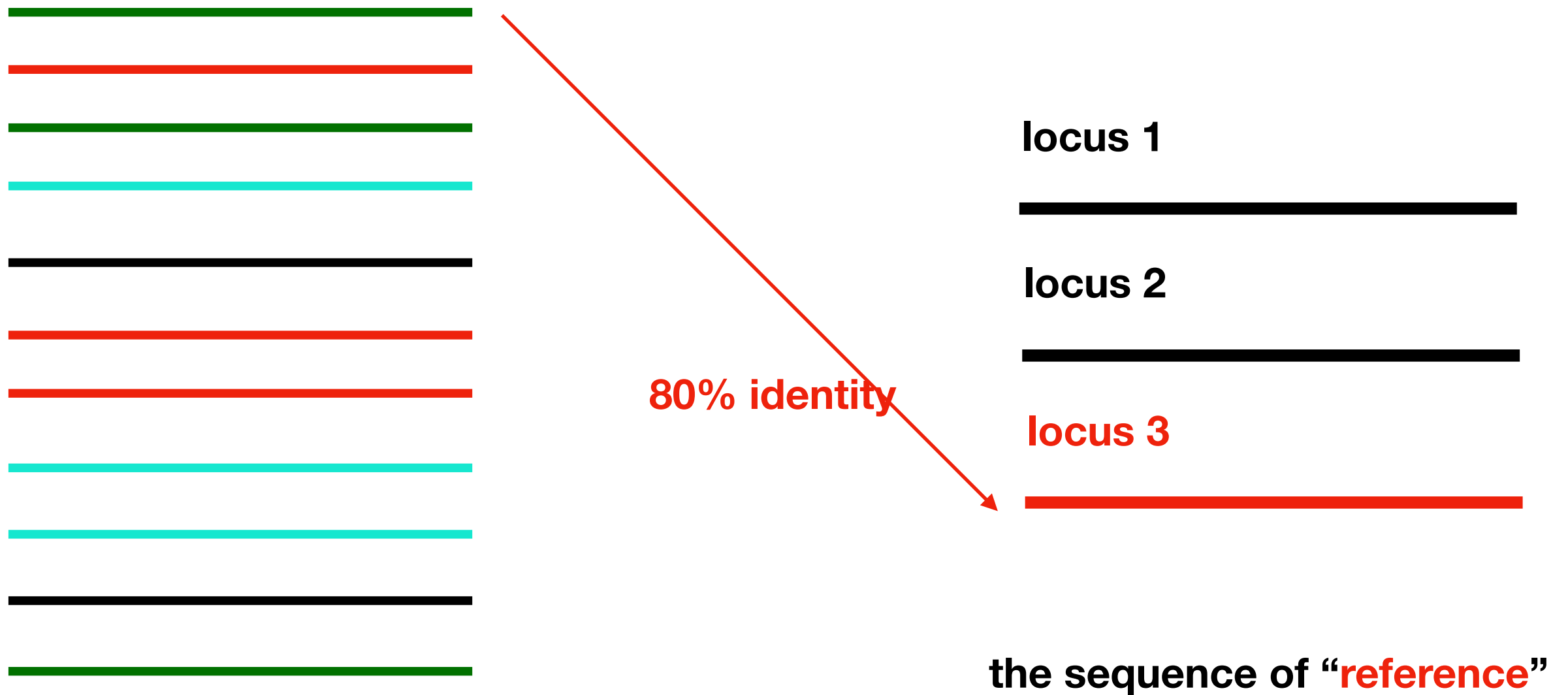
Parse reads to loci



Mixed short reads from several loci

Reads got different identity with each locus

Parse reads to loci



Mixed short reads from several loci

Select reads with highest identity

Parse reads to loci

locus 1

locus 2

locus 3

Unassigned reads

Most of reads are assigned to different loci.
Some reads from nowhere are still unassigned

Assemble parsed reads

locus 1

—————

—————

—————

—————



—————

locus 2

—————

—————

—————

—————



—————

locus 3

—————

—————

—————

—————



—————

Assemble parsed reads into longer contigs

Assemble parsed reads

Reads:



Contig:



Assemble parsed reads



In real case, question is not that easy. We always have loci assigned with more than 2,000 reads

Assemble parsed reads

Find overlaps among all reads

Build a graph recording overlaps among all reads

Traverse through the graph to get contigs

Assemble parsed reads

Find overlaps among all reads

Align the reads

K-mer

FM-index

Assemble parsed reads

Find overlaps among all reads

Align the reads

K-mer

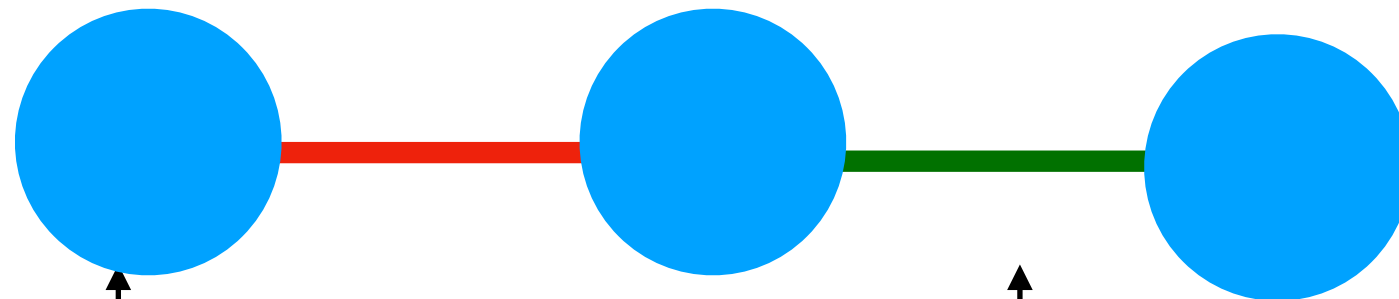
FM-index

Only considerable length of overlap between reads will be kept (25 bp), to guarantee the low probability of accidentally overlap occurs

Assemble parsed reads

Build a graph recording overlaps among all reads

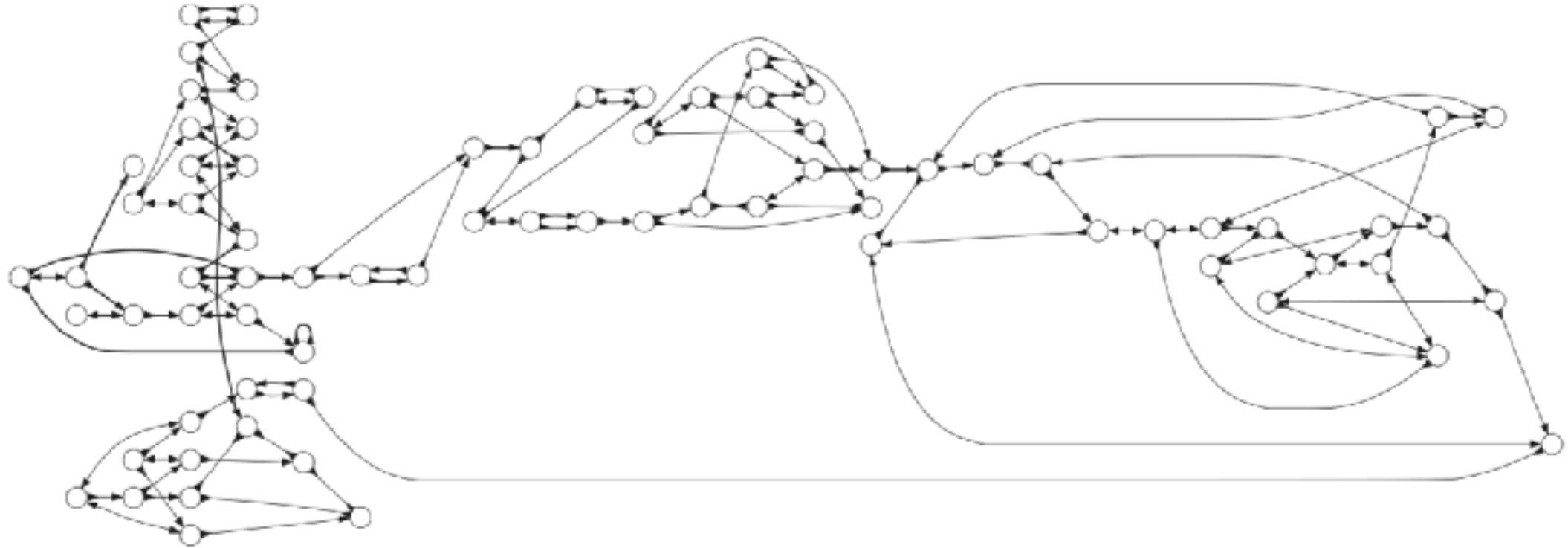
Reads:



Circle represent
the read we call
it “**node**” or
“**vertex**”

Line represents the
connection
between reads
called “**edge**”

Assemble parsed reads



Graph of *Campylobacter jejuni*

Assemble parsed reads

Traverse through the graph to get contigs

For each locus, there's only one sequence

Traverse through the all nodes in the graph and each node only pass once

Assemble parsed reads

Traverse through the graph to get contigs

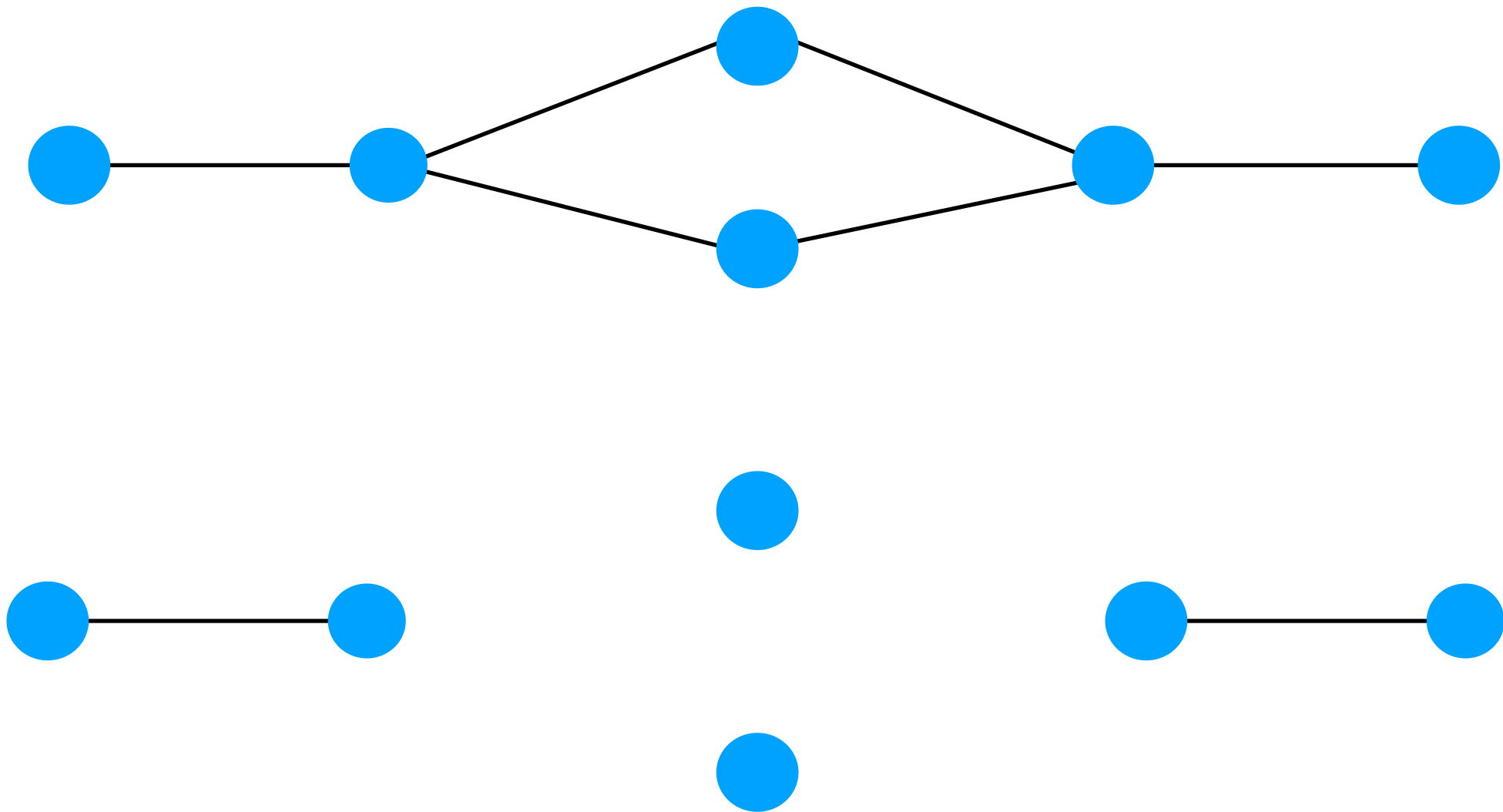
For each locus, there's only one sequence

Traverse through the all nodes in the graph and each node only pass once

But this assumption is hard to fulfill

Assemble parsed reads

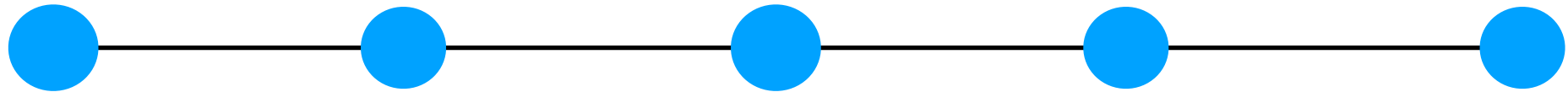
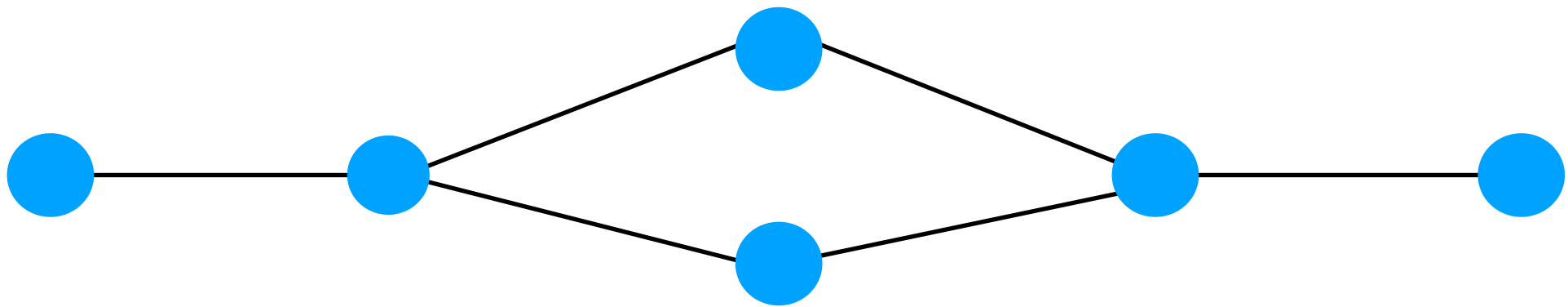
Bubble:



Break the graph into several sub-graph

Assemble parsed reads

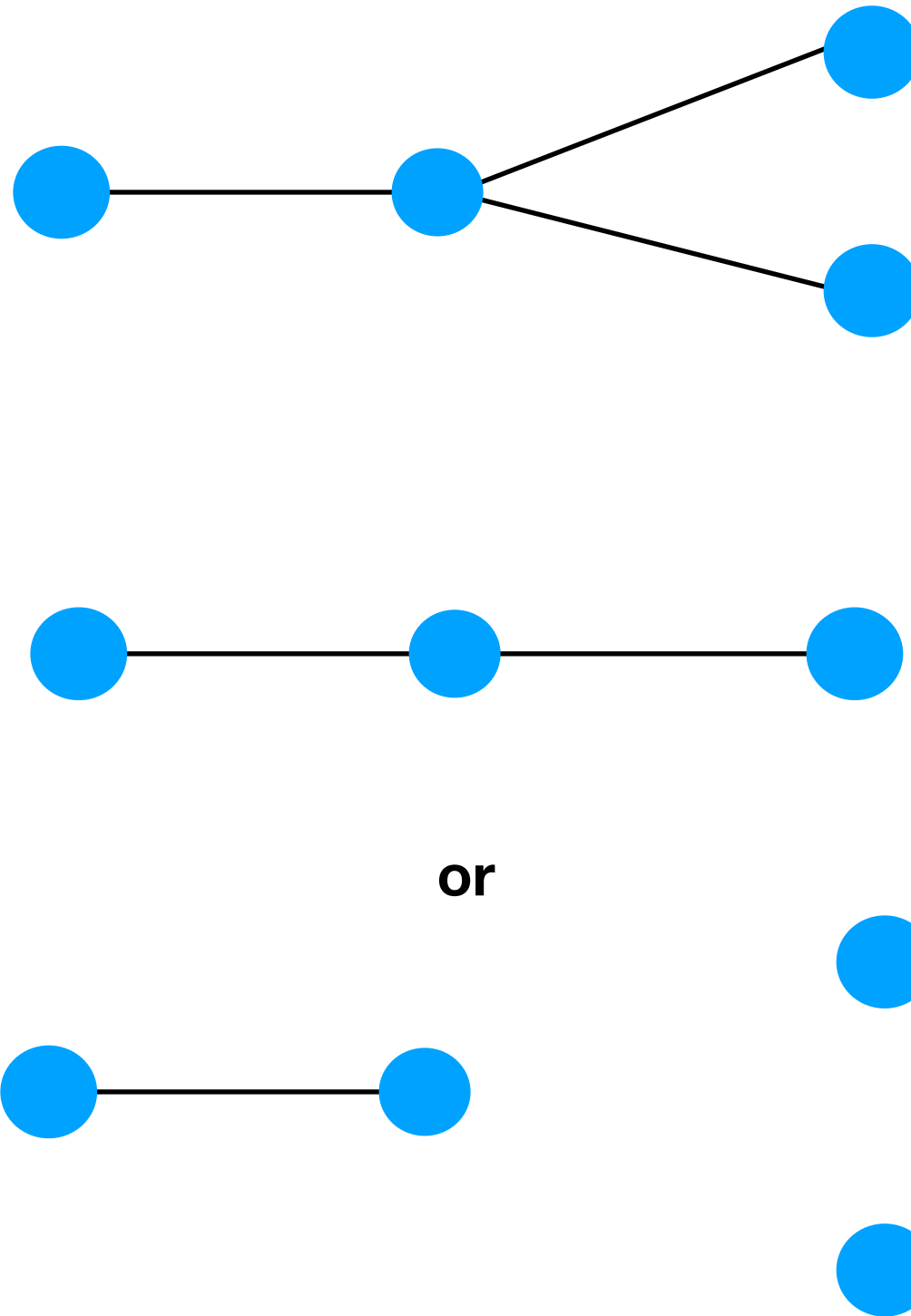
Bubble:



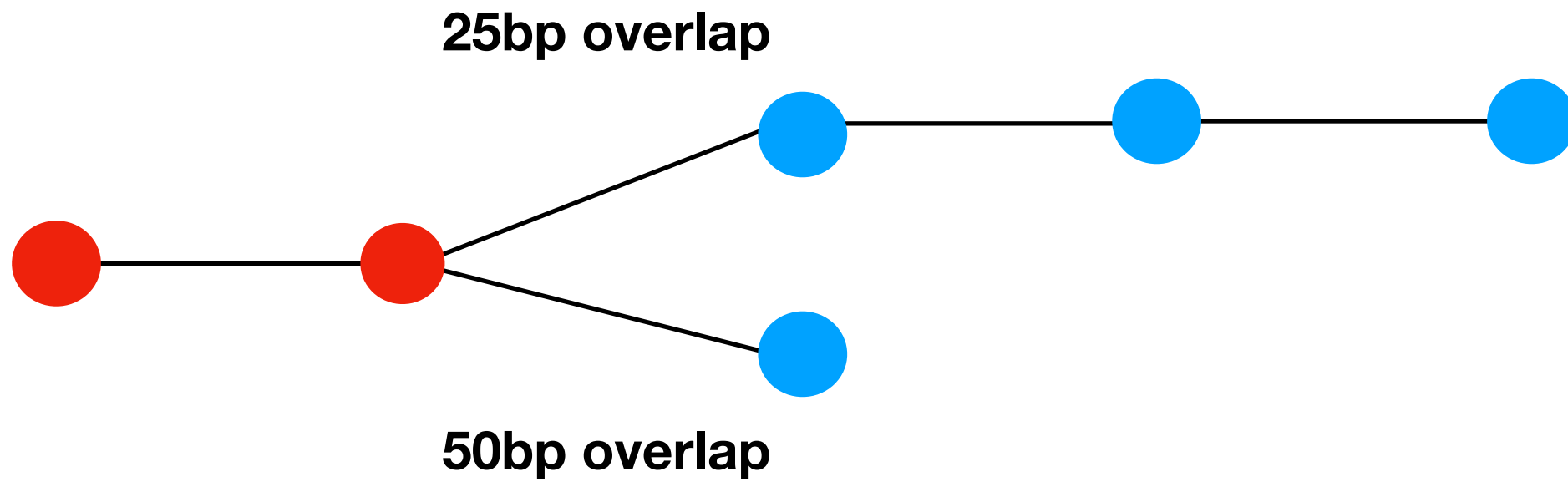
Discard one of the path

Assemble parsed reads

Tip:



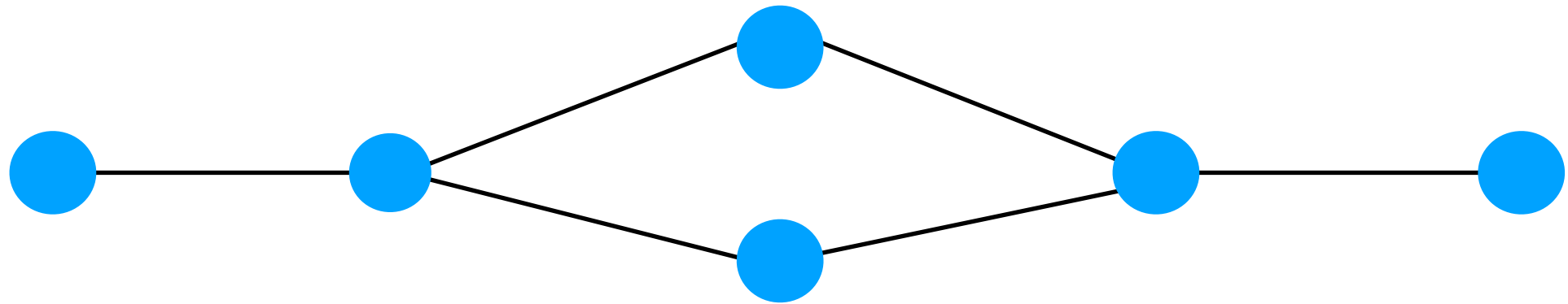
Why we need graph



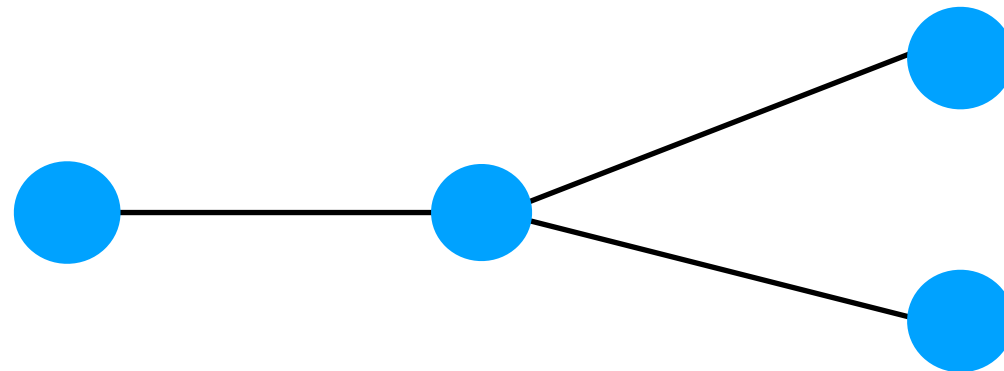
Which way you should choose

Further assemble

Bubble:



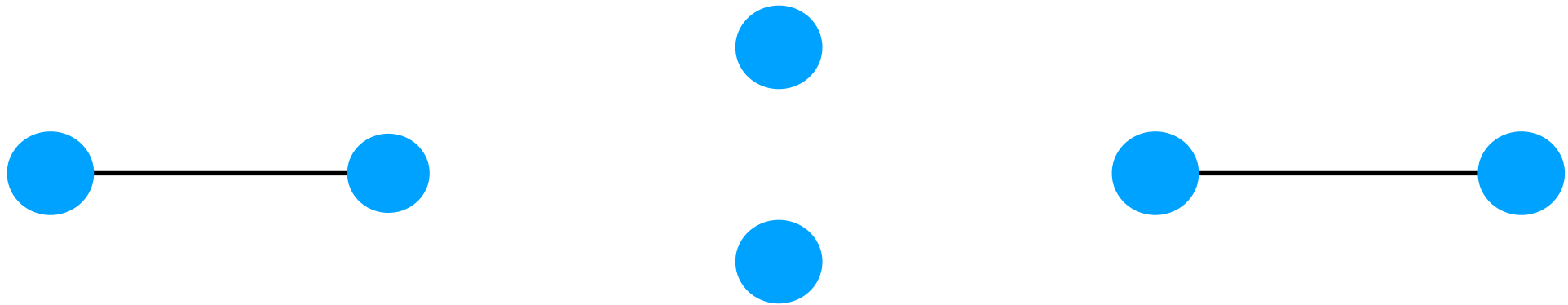
Tip:



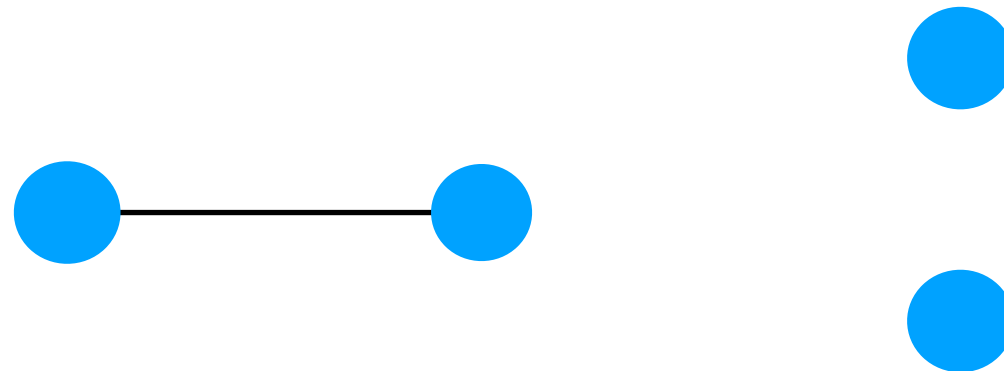
If two path is too diverged ($\geq 95\%$ identity). The path will be split into several contigs

Further assemble

Bubble:



Tip:



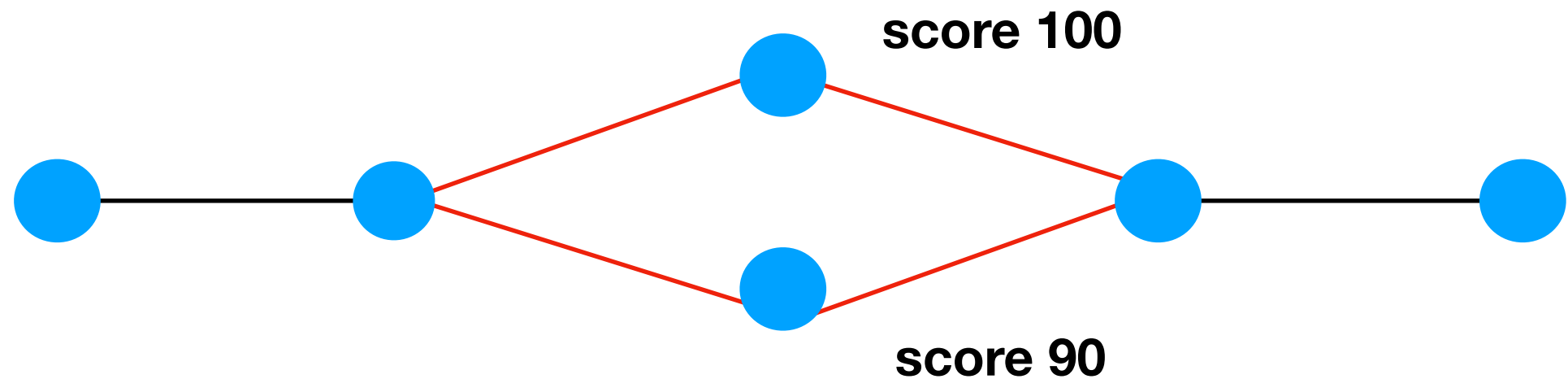
If two path is too diverged ($\geq 95\%$ identity). The path will be split into several contigs

Further assemble

locus1:



Bubble:



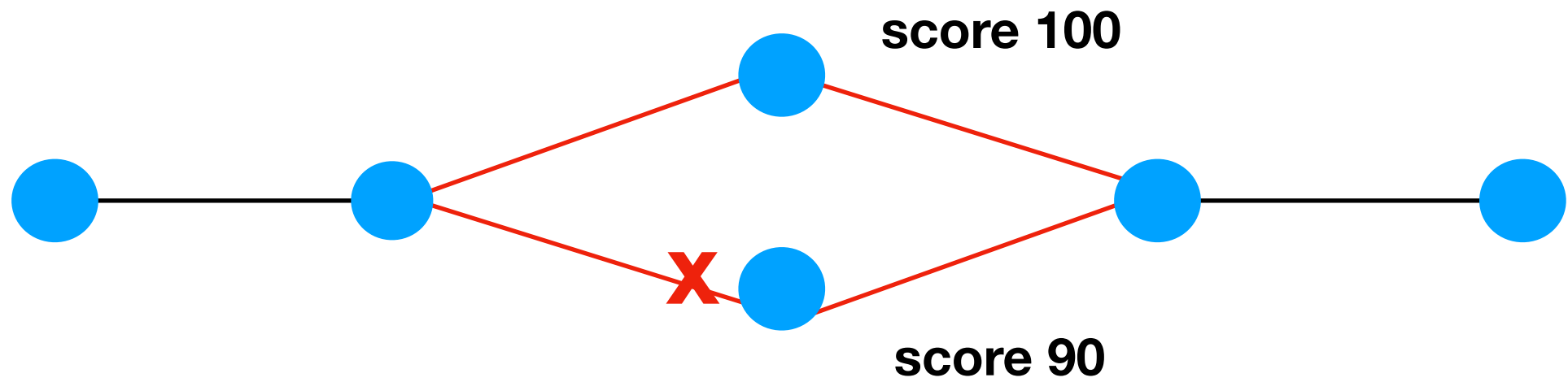
Each contig will be aligned to reference. Graph will be reconstructed. The alignment score of each path will be calculated.

Further assemble

locus1:



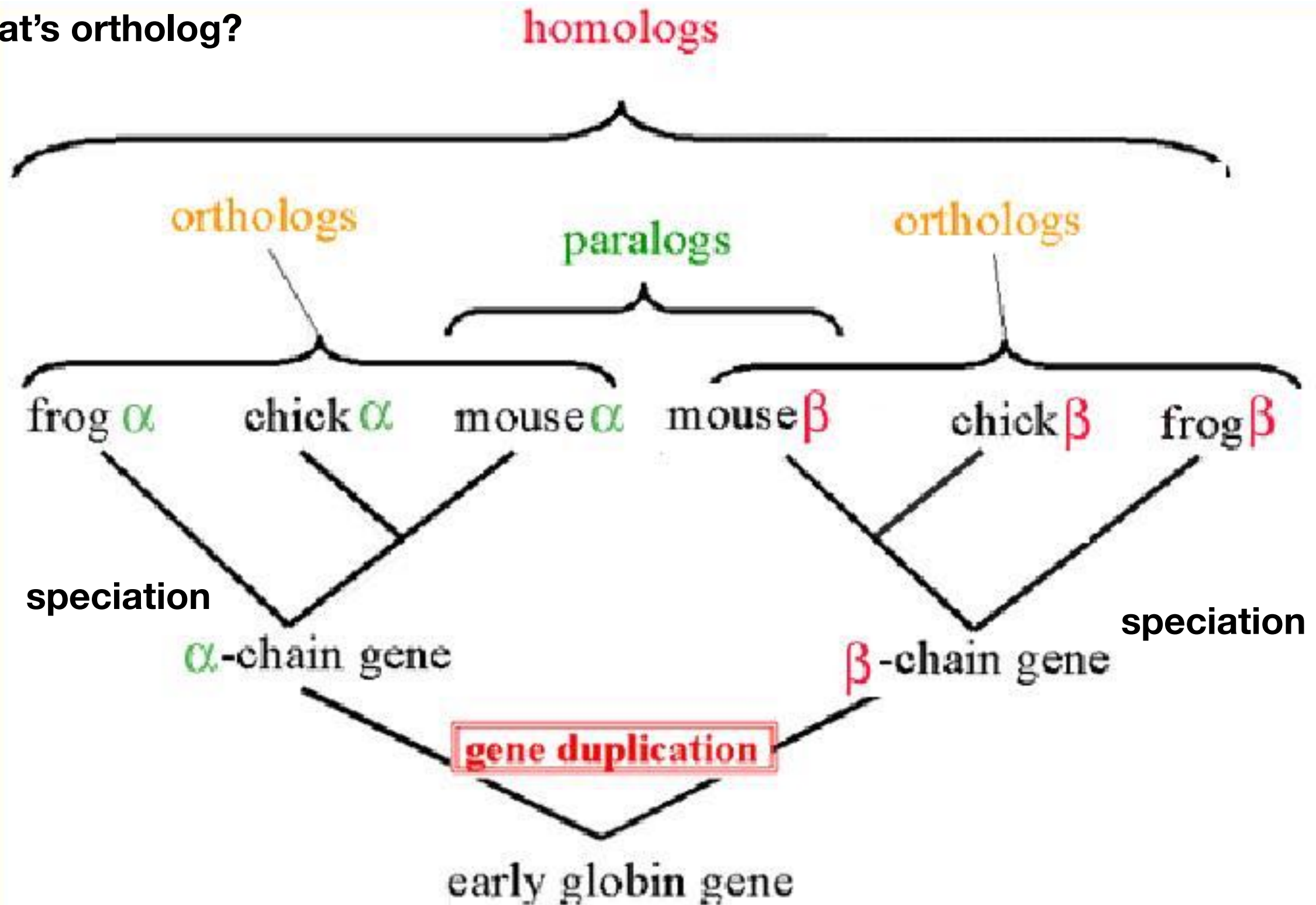
Bubble:



Keep the path with higher score.

Get orthologous assemblies

What's ortholog?



Get orthologous assemblies

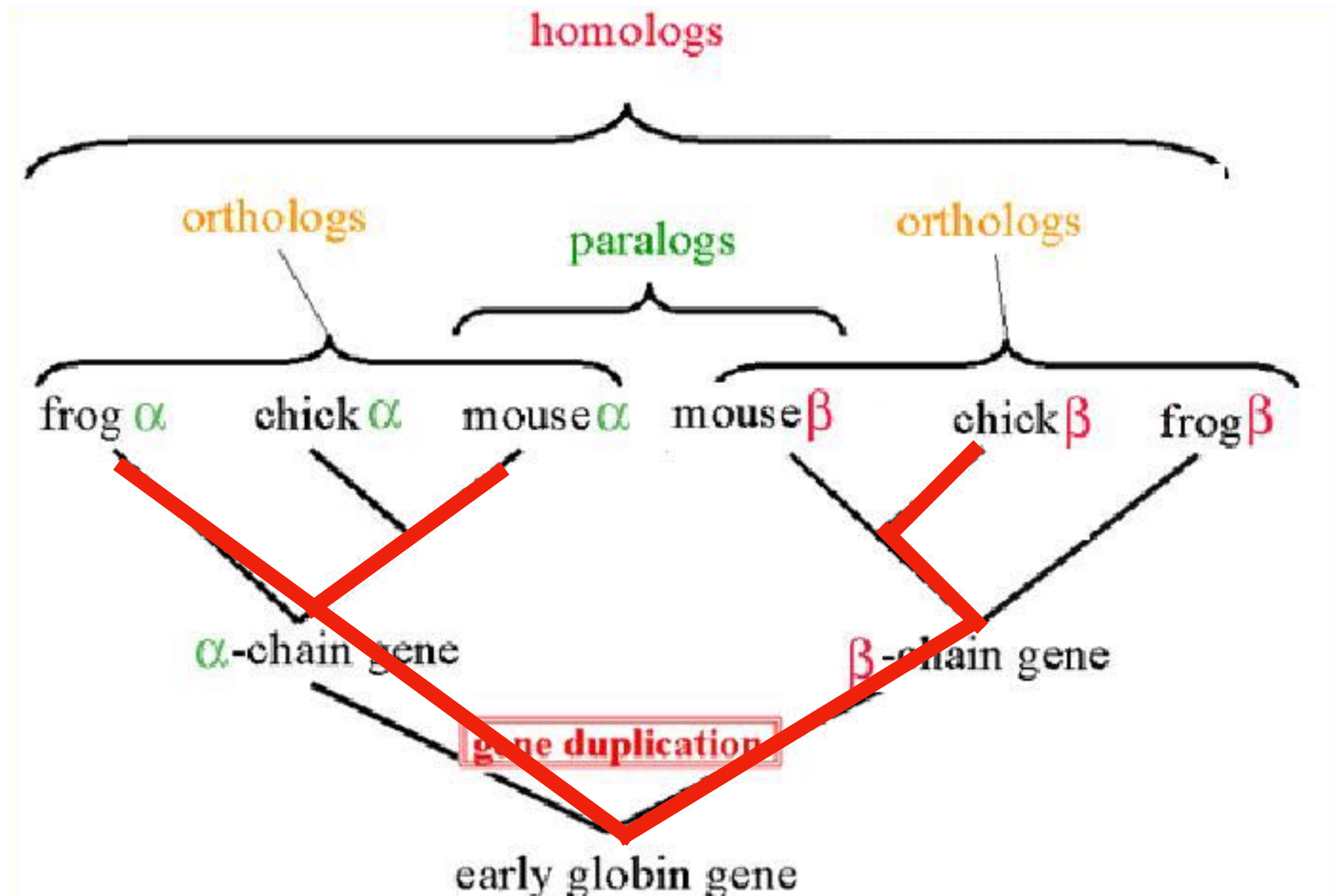
Why we need it?

Our final aim is to reveal evolutionary history among our enriched species

Orthologues genes are derived from “**speciation event**”. So, the evolutionary history of these genes are identical with the evolutionary history of species

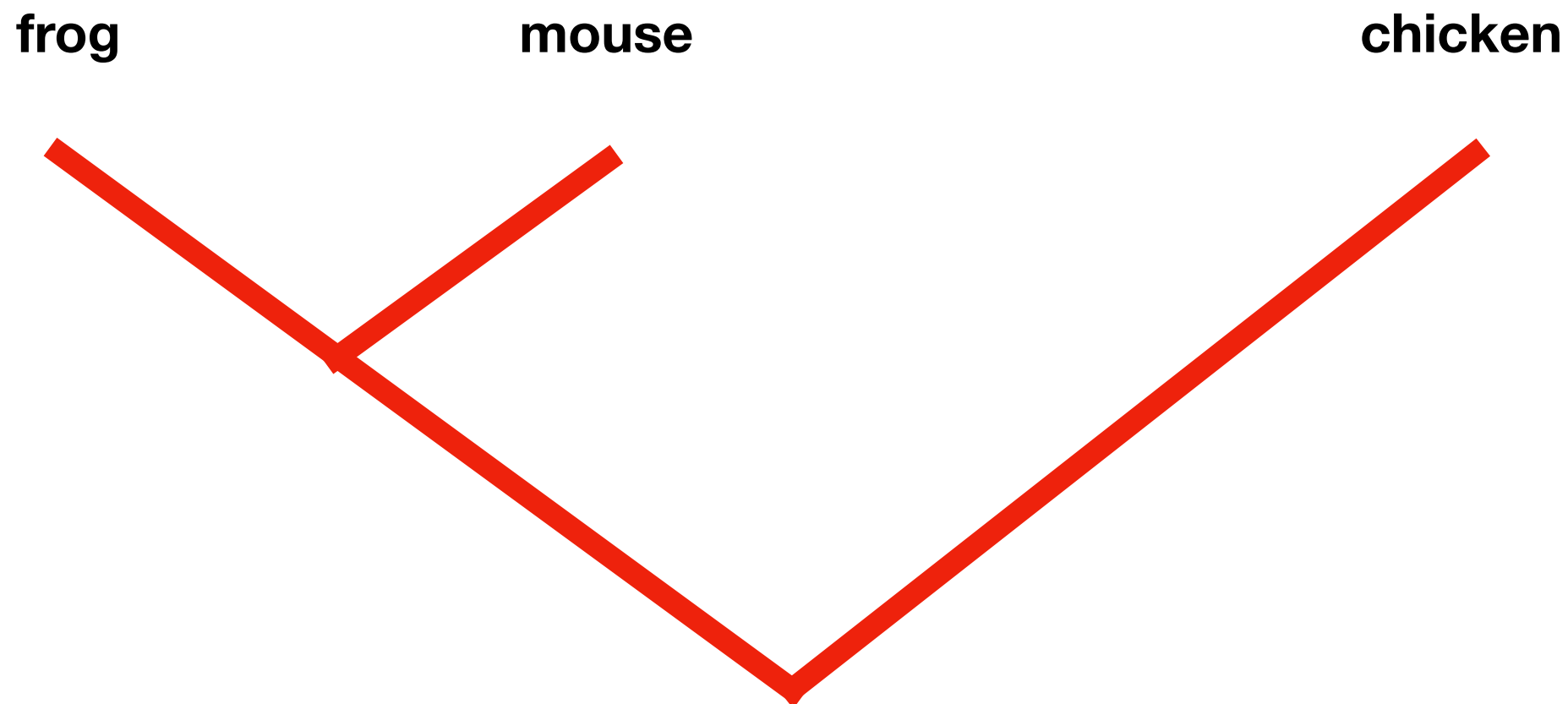
Get orthologous assemblies

What will happen if we use paralog genes to reveal evolutionary history



Get orthologous assemblies

What will happen if we use paralog genes to reveal evolutionary history



Get orthologous assemblies

There's various way to find orthologues. The method we used here called **reciprocal blast**

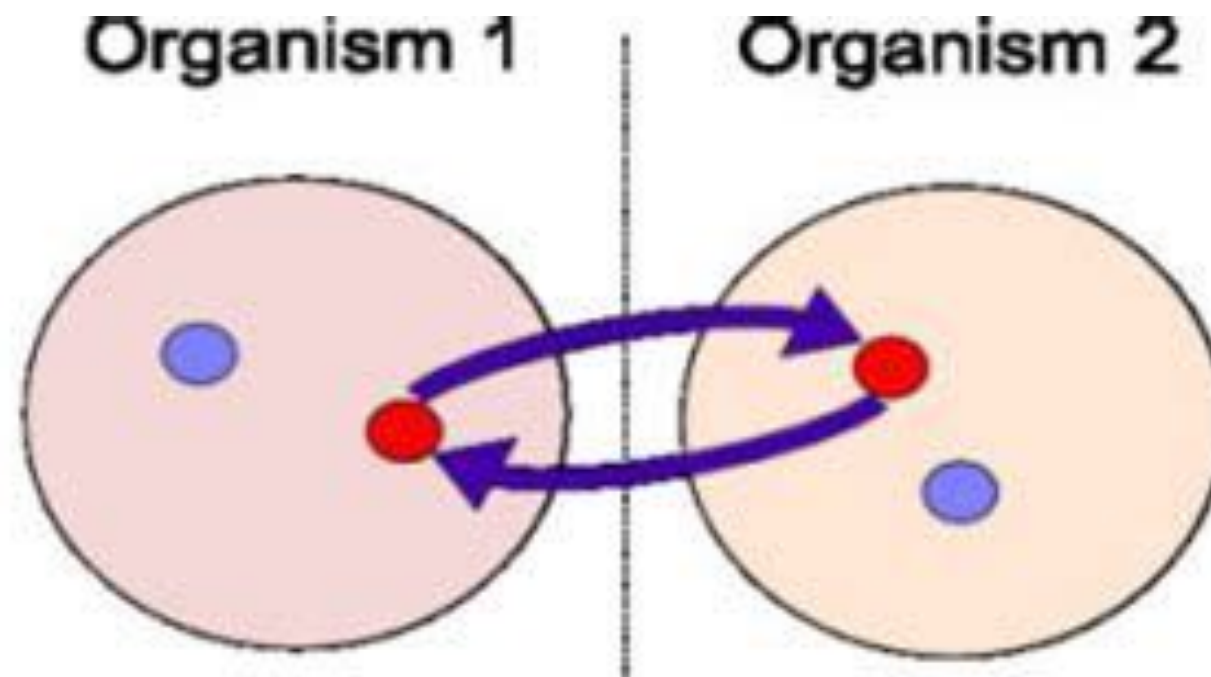
This method is built on the assumption that **orthologous genes have identical or highly related functions and this sharing is greater than for paralogs.**

Closest gene between 2 species are potential orthologous gene

Get orthologous assemblies

Reciprocal blast in general case

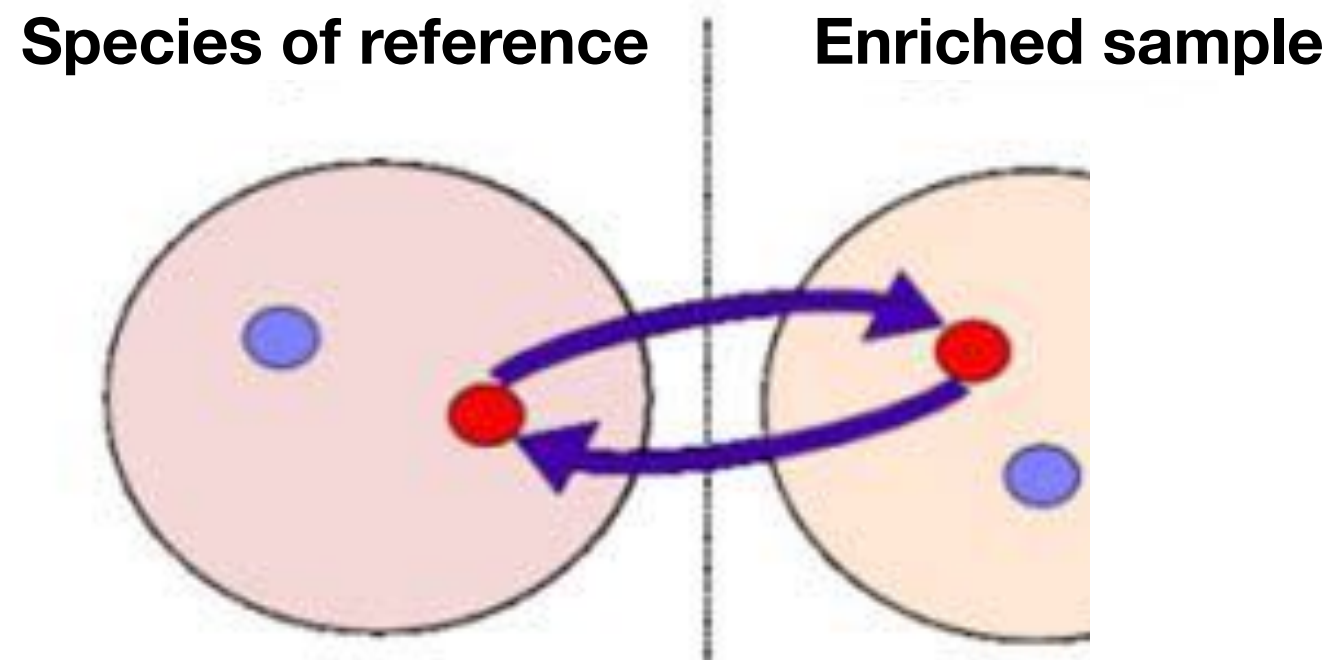
- (1) sequence of gene which we want to find its orthologous sequence in other organisms
- (2) genome of these 2 organisms



If a pair of genes in different species are the closest to each other, these 2 genes have “**reciprocal best hit**”

Get orthologous assemblies

Reciprocal blast in our pipeline

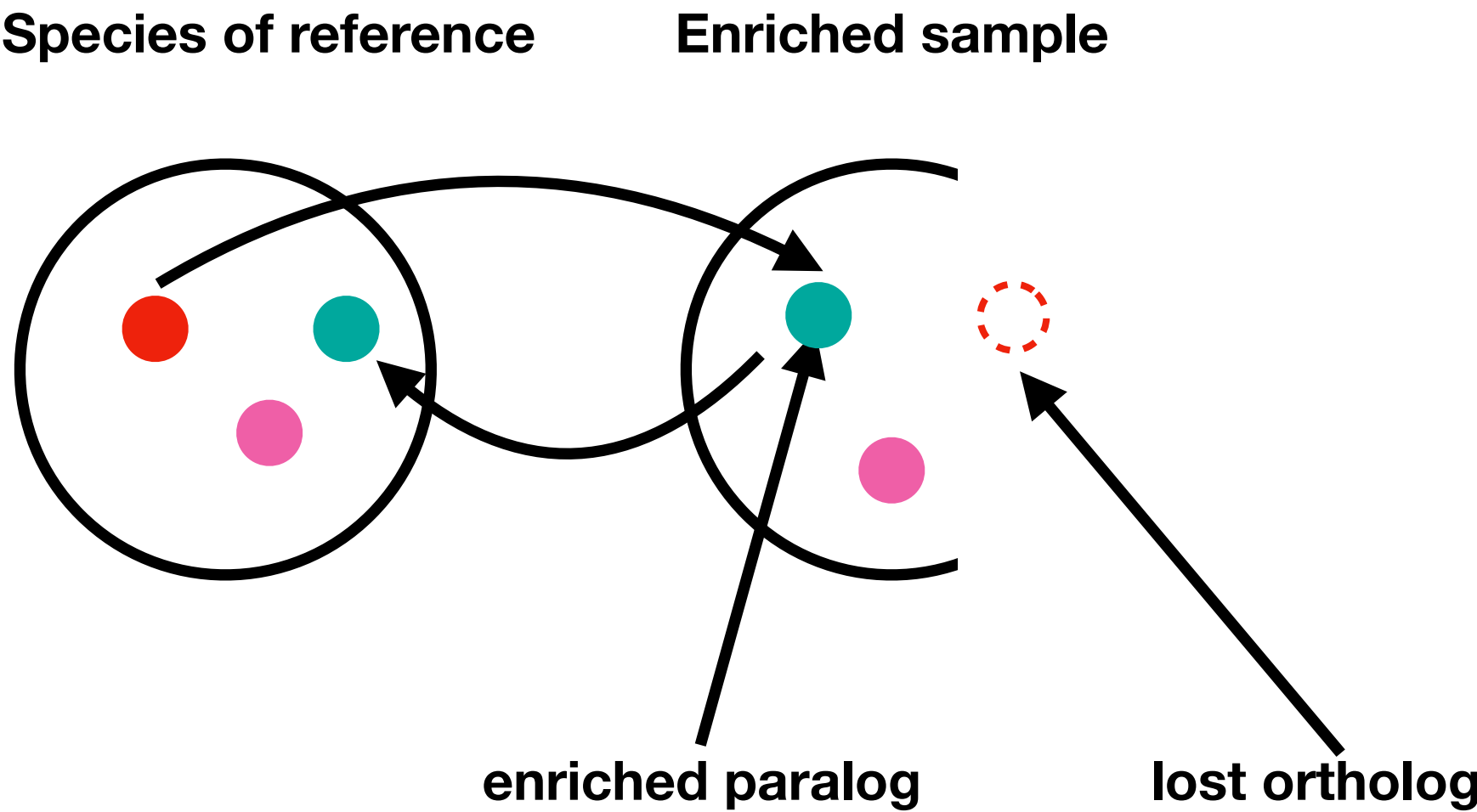


In our situation, we do not have the genome, only got several contigs only

But, loci of reference and contig have reciprocal best hit, then they are also putative orthologues

Get orthologous assemblies

Reciprocal blast in our pipeline



Exclude the contig if it does not have reciprocal blast hit with reference loci

Until here we've been reached our first goal

