

Workshop on Genomics 2015

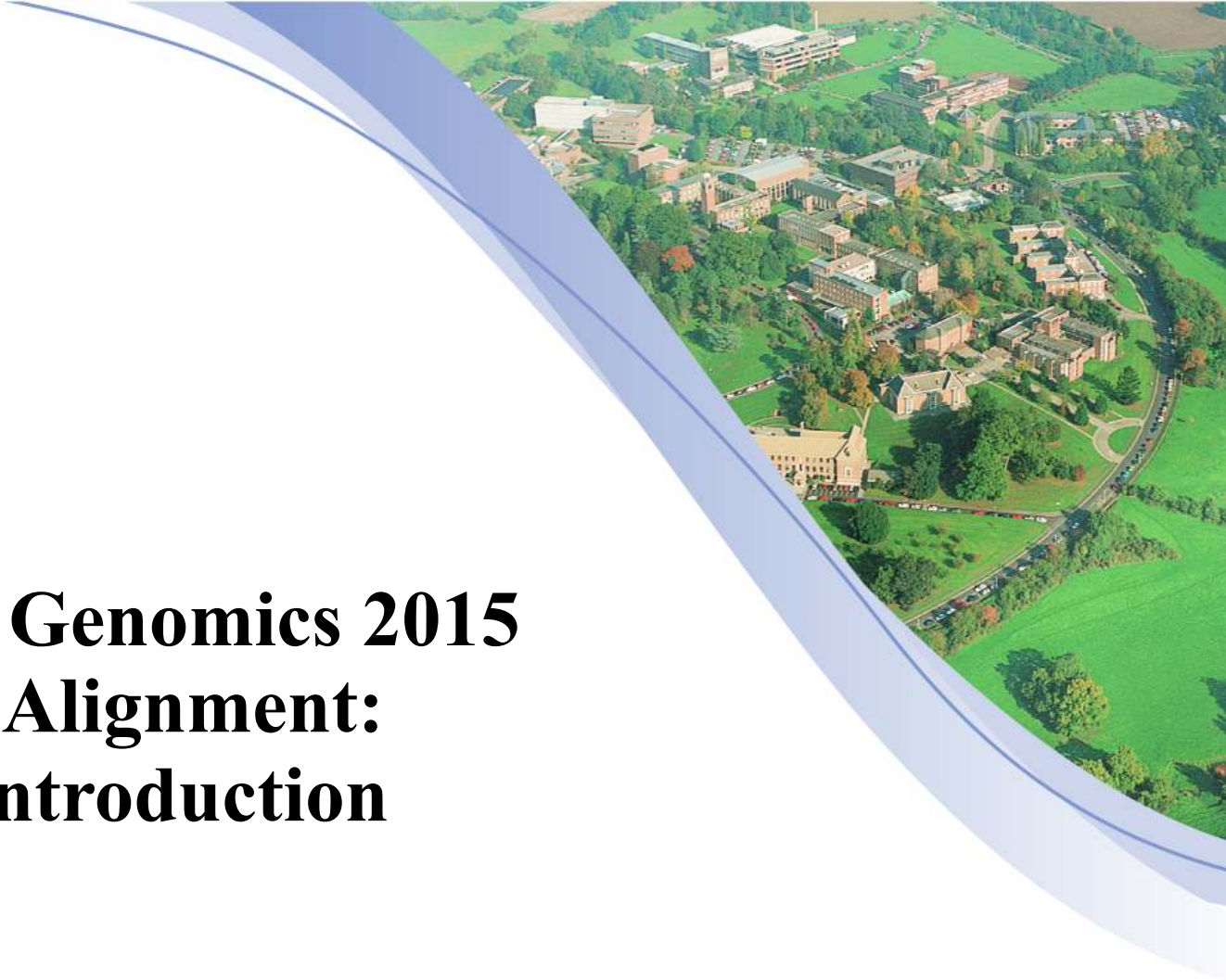
Sequence Alignment:

An brief introduction

Konrad Paszkiewicz

University of Exeter, UK

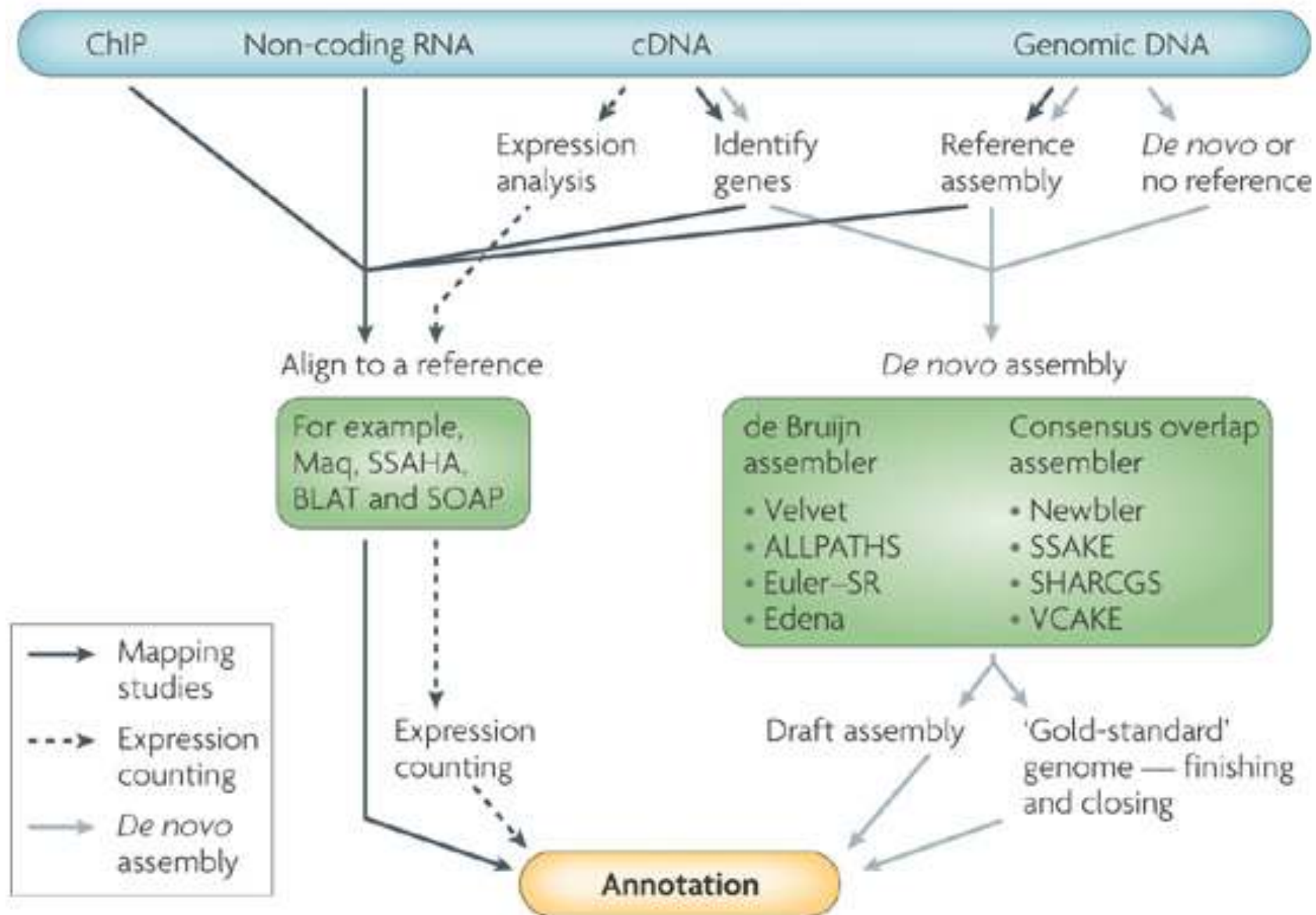
k.h.paszkiewicz@exeter.ac.uk



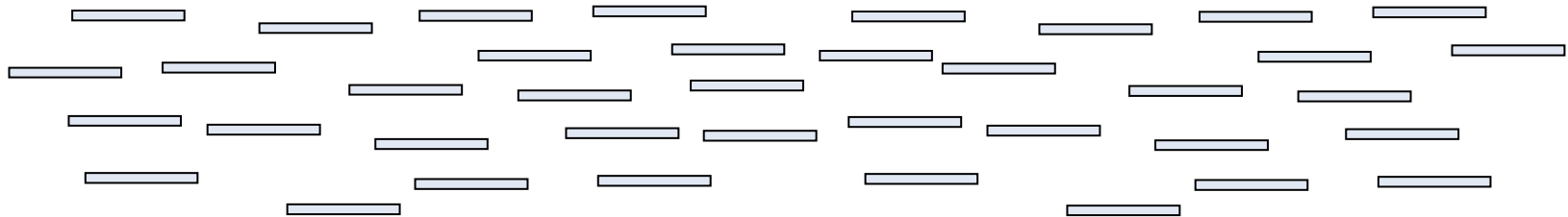
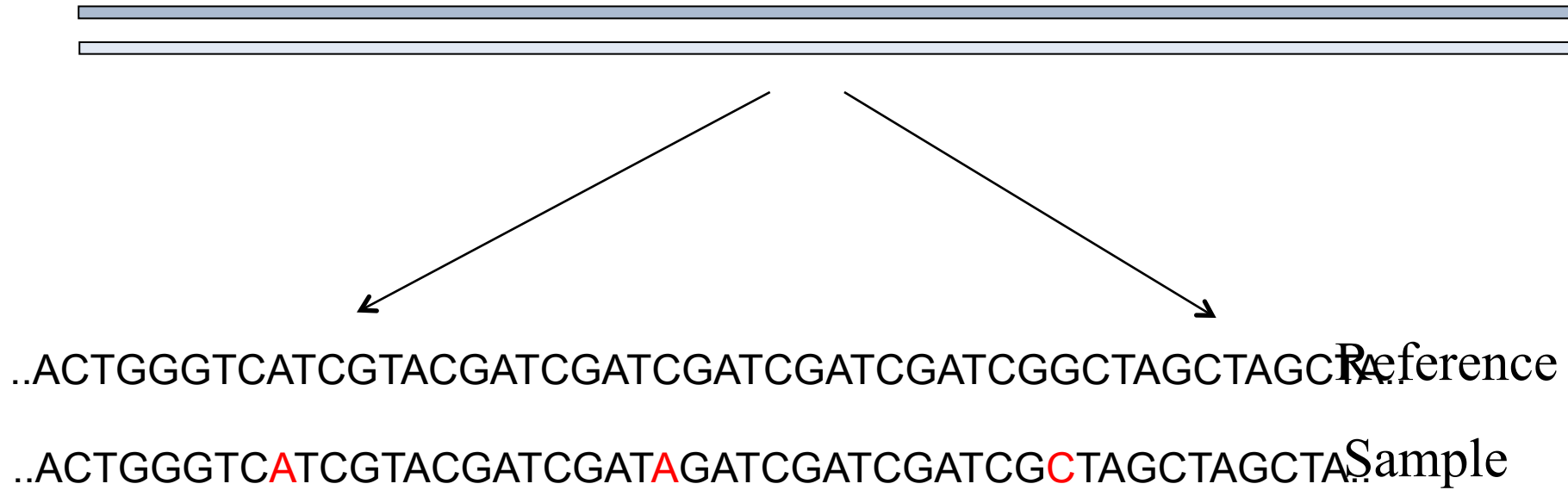
Contents

- **Alignment algorithms for short-reads**
 - Background – Blast (why can't we use it?)
 - Adapting hashed seed-extend algorithms to work with shorter reads
 - Suffix/Prefix Tries
 - Indels
 - Other alignment considerations
 - Typical alignment pipeline
 - Haplotype methods of variant calling

Raw sequence source



Alignment of reads to a reference



Why is short read alignment hard?

The shorter a read, the less likely it is to have a unique match to a reference sequence

Need over 7kb reads to uniquely place most reads from bacteria

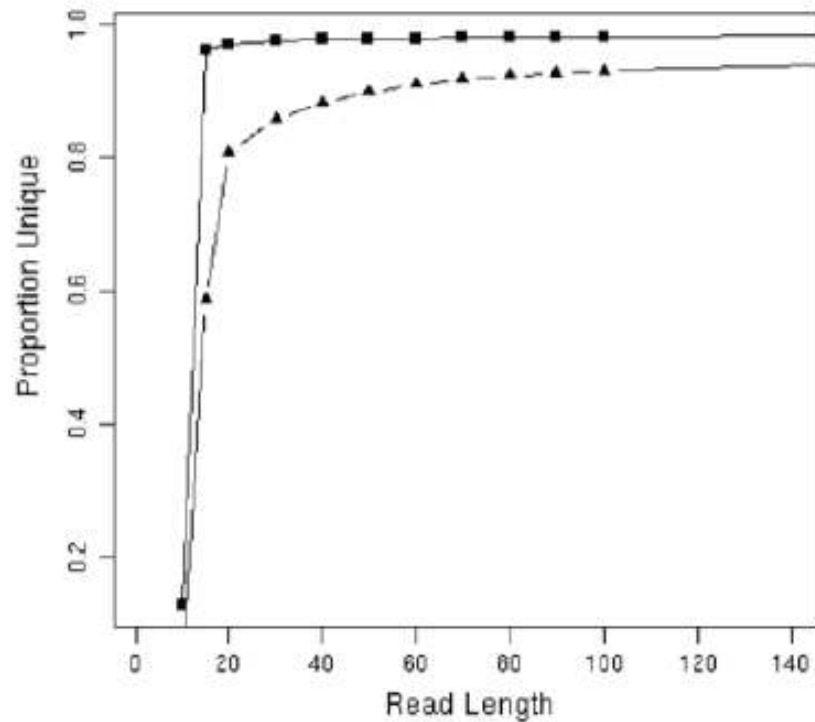


Fig. 1 The proportion of unique sequence in the *Streptococcus suis* (squares) and *Mus musculus* (triangles) genomes for varying read lengths. This graph indicates that read length has a critical affect on the ability to place reads uniquely to the genome

Why do we generate short reads?

- Sanger reads lengths $\sim$ 800-2000bp
- Generally we define short reads as anything below this
 - Illumina (50bp – 300bp)
 - SoLID (80bp max)
 - Ion Torrent (200-400bp max...)
 - Roche 454 – 400-800bp
- Even with these platforms it is cheaper to produce short reads (e.g. 50bp) rather than 100 or 200bp reads
- Diminishing returns:
 - For some applications 50bp is more than sufficient
 - Small-RNA
 - ChIP-Seq
 - Differential gene expression
 - Digital Gene Expression profiling

Short read alignment applications

Genotyping:

Methylation

SNPs

Indels

```
...CCATAG      TATGCGCCC      CGGAATTTCGGTATAC...
...CCAT      CTATATGCG      TCGGAATTTCGGTATAC
...CCAT  GGCTATATG      CTATCGGAAA  GCGGTATA
...CCA  AGGCTATAT      CCTATCGGA      TTGCGGTAC  C...
...CCA  AGGCTATAT      GCCCTATCG      TTTGCGGT  C...
...CC   AGGCTATAT      GCCCTATCG      AAATTTGC  ATAC...
...CC   TAGGCTATA  GCGCCCTA      AAATTTGC  GTATAC...
...CCATAGGCTATATGCGCCCTATCGGCAATTTGCGGTATAC...
```

Classify and measure peaks:

ChIP-Seq

RNA-Seq

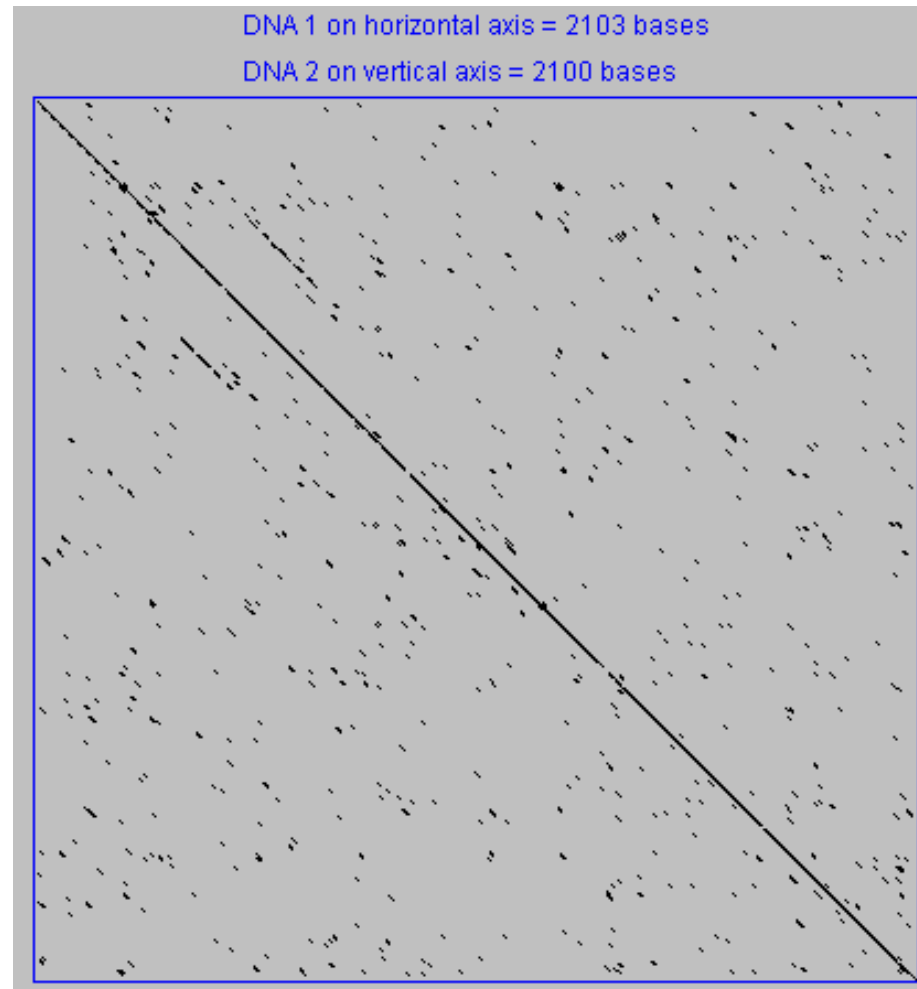
```
...CC
...CCATAGGCTATATGCGCCCTATCGGCAATTTGCGGTATAC...
GAAATTTGC
GGAAATTTG
CGGAAATTT
CGGAAATTT
TCGGAAATT
CTATCGGAAA
CCTATCGGA  TTTGCGGT
GCCCTATCG  AAATTTGC
GCCCTATCG  AAATTTGC  ATAC...
```

Contents

- **Alignment algorithms for short-reads**
 - **Background – Blast (why can't we use it?)**
 - Global alignment
 - Local alignment
 - Adapting hashed seed-extend algorithms to work with shorter reads
 - Indel detection
 - Suffix/Prefix Tries
 - Other alignment considerations
 - Typical alignment pipeline
 - Haplotype methods of variant calling

Dot Matrix Method

- Aligning by eye



Sequence Alignment

ATCGATA-CG

AT**G**GAT**T**ACG

3 possibilities

Match

...A...
|
...A...

Mismatch

...**C**...
...**G**...

Indel

...-...
...**T**...

A very simple alignment scoring system

Points for a matching letter: 1

Points for a non-matching letter: 0

Points for inserting a gap: 0

Global Pair-wise Alignment

ATCGATACG, ATGGATTACG

ATCGAT-ACG
| | | | |
ATGGATTACG

Matches:	+1	+1		+1	+1	+1		+1	+1	+1	= +8
Mismatches:			0								= 0
Gaps:							0				= 0
											Total score = +8

But, what does this score mean??
Could we get a better alignment?

How to choose the best alignment?

- Sequence 1: ACTGAGC
- Sequence 2: ATGATGC
- Some possible alignments:

ACTGAGC-- ACTGA-GC A-----CTGAGC
A-TGA-TGC A-TGATGC ATGAT-----GC

Global alignment – Needleman-Wunsch

A **global** alignment covers the **entire lengths of the sequences** involved

The Needleman-Wunsch algorithm finds the best global alignment between 2 sequences across their whole length

Step 1: Initialise

	A	C	T	G	A	G	C
A							0
T							0
G							0
A							0
T							0
G							0
C	0	1	0	0	0	0	1

Fill in far-right column and bottom row with:

0 for a mis-match

1 for a match

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A							0
T							0
G							0
C	0	1	0	0	0	0	1

For each box, find the highest number out of the blue boxes

Step 3:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A							0
T							0
G						1+1=2	0
C	0	1	0	0	0	0	1

If there is a match in the yellow box as, take the highest value from the blue boxes and add 1 to it

G matches G in the yellow box, so add 1 to the 1 in the blue box

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A							0
T							0
G					0+1=1	2	0
C	0	1	0	0	0	0	1

A does not match G. So add zero to the 1 in the blue box.

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A							0
T							0
G				1+1=2	1	2	0
C	0	1	0	0	0	0	1

If there is a match as here, take the highest value and add 1 to it

G matches G so add 1 to 1 in the blue boxes

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A							0
T							0
G			0+1=1	2	1	2	0
C	0	1	0	0	0	0	1

If there is a match as here, take the highest value and add 1 to it

T does not match G. So add zero.

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A							0
T						0+1=1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Highest out of the blue boxes is 1

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A							0
T					2+0=2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Highest out of the blue boxes is 2

A does not match T

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A							0
T				2+0=2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Highest out of the blue boxes is 2

G does not match T

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A							0
T			3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Highest out of the blue boxes is 2

T does match T

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A							0
T		2+0=2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Highest out of the blue boxes is 2

C does not match T

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A							0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Do the same for all remaining rows

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A						1+0=1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Do the same for all remaining rows

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A					2+1=3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Do the same for all remaining rows

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A				2+0=2	3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Do the same for all remaining rows

Step 2:

	A	C	T	G	A	G	C
A							0
T							0
G							0
A			2+0=2	2	3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Do the same for all remaining rows

Step 2:

	A	C	T	G	A	G	C
A	6	5	4	3	3	1	0
T	4	4	5	3	2	1	0
G	3	3	3	4	2	1	0
A	4	3	2	2	3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Do the same for all remaining rows

Step 3: Backtracking

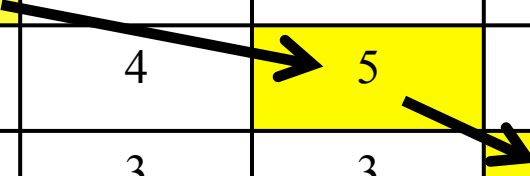
	A	C	T	G	A	G	C
A	6	5	4	3	3	1	0
T	4	4	5	3	2	1	0
G	3	3	3	4	2	1	0
A	4	3	2	2	3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1



Follow largest numbers starting from top-left going down and to the right

Step 3: Backtracking

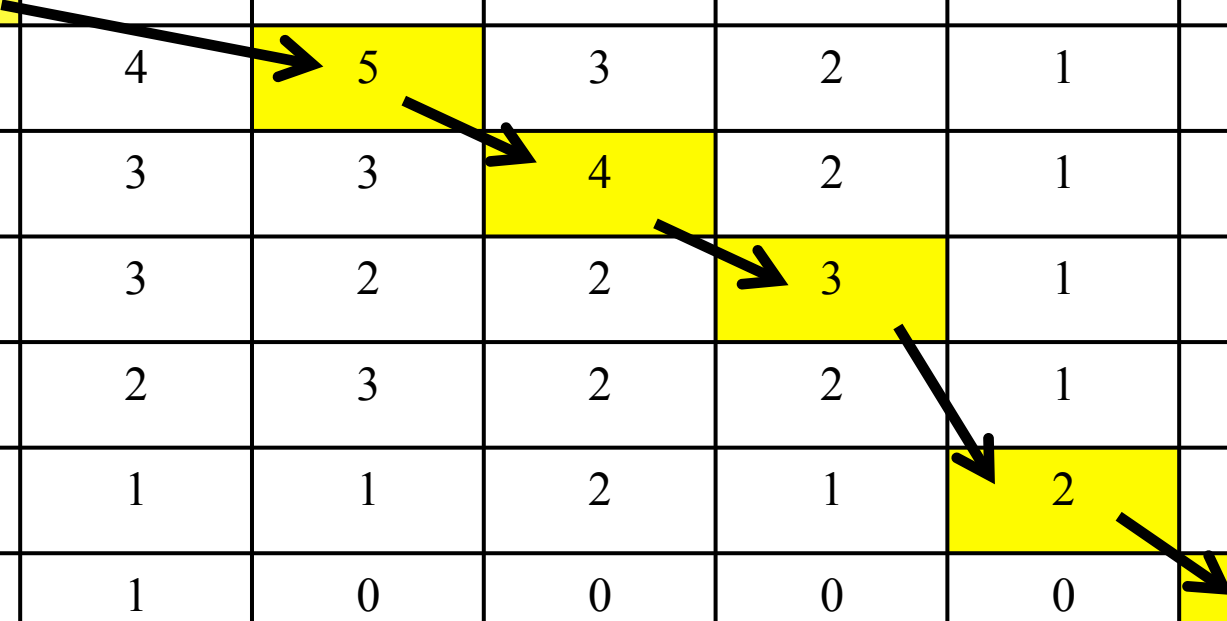
	A	C	T	G	A	G	C
A	6	5	4	3	3	1	0
T	4	4	5	3	2	1	0
G	3	3	3	4	2	1	0
A	4	3	2	2	3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1



Follow largest numbers starting from top-left going down and to the right

Step 3: Backtracking

	A	C	T	G	A	G	C
A	6	5	4	3	3	1	0
T	4	4	5	3	2	1	0
G	3	3	3	4	2	1	0
A	4	3	2	2	3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1



Follow **largest** numbers starting from **top-left**,
going **down** and to the **right**

Step 4: Generate alignment

	A	C	T	G	A	G	C
A	6	5	4	3	3	1	0
T	4	4	5	3	2	1	0
G	3	3	3	4	2	1	0
A	4	3	2	2	3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Horizontal seq

A

Vertical seq

A

Step 4: Generate alignment

Gap

	A	C	T	G	A	G	C
A	6	5	4	3	3	1	0
T	4	4	5	3	2	1	0
G	3	3	3	4	2	1	0
A	4	3	2	2	3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Horizontal seq
Vertical seq

ACT
A-T

Step 4: Generate alignment

	A	C	T	G	A	G	C
A	6	5	4	3	3	1	0
T	4	4	5	3	2	1	0
G	3	3	3	4	2	1	0
A	4	3	2	2	3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Horizontal seq
Vertical seq

ACTG
A-TG

Step 4: Generate alignment

	A	C	T	G	A	G	C
A	6	5	4	3	3	1	0
T	4	4	5	3	2	1	0
G	3	3	3	4	2	1	0
A	4	3	2	2	3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Horizontal seq
Vertical seq

ACTGA
A-TGA

Step 4: Generate alignment

	A	C	T	G	A	G	C
A	6	5	4	3	3	1	0
T	4	4	5	3	2	1	0
G	3	3	3	4	2	1	0
A	4	3	2	2	3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Horizontal seq
Vertical seq

ACTGA-
A-TGAG

Step 4: Generate alignment

	A	C	T	G	A	G	C
A	6	5	4	3	3	1	0
T	4	4	5	3	2	1	0
G	3	3	3	4	2	1	0
A	4	3	2	2	3	1	0
T	2	2	3	2	2	1	0
G	1	1	1	2	1	2	0
C	0	1	0	0	0	0	1

Horizontal seq
Vertical seq

ACTGA-C
A-TGAGC

Optimal global alignment

ACTGA-C

| | | | |

A-TGAGC

Local alignment

A **global** alignment is often not appropriate as only parts of sequences may be conserved

A **local** alignment only covers **parts of the sequences**

The **Smith-Waterman** algorithm finds the **best local alignment** between 2 sequences

Global alignment

		Q	K	E	S	G	P	S	S	S	Y	C
	V	Q	Q	E	S	G	L	V	R	T	T	C

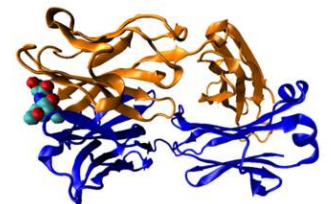
Local alignment

		E	S	G
		E	S	G

Local alignment

A **local alignment** of 2 sequences is an alignment between **parts** of the 2 sequences

E.g. Two proteins may be very similar in a functional site, but be very dissimilar outside that region



A global alignment of such sequences would have:

- (i) lots of matches in the region of high sequence similarity
- (ii) lots of mismatches & gaps (insertions/deletions) outside the region of similarity

It makes sense to find the **best local alignment** instead

```

human/1-422      . . . . .
fly/1-898      1 MFTLQPTPTAIGTVPPWSAGTLIERLP SLEDMAHKDNV I AMRNLPCLGT 50

human/1-422      1 . . . . .MQNSHSG 7
fly/1-898      51 AGGSGGLGG I AGKPSPTMEAVEASTASHPHSTSSYFATTYYHLTDDECHSG 100

human/1-422      8 VNQLGGVFVN GRPLPDSTRQKIVELAHSGARPCDISRILQVSNGCVSKIL 57
fly/1-898      101 VNQLGGVFVG GRPLPDSTRQKIVELAHSGARPCDISRILQVSNGCVSKIL 150

human/1-422      58 GRYYETGSIRPRAIGGSKPRVATPEVVSKI AQYKRECPSIF AWEIRDRL 107
fly/1-898      151 GRYYETGSIRPRAIGGSKPRVATAEVVSKI SQYKRECPSIF AWEIRDRL 200

human/1-422      108 SEG VCTNDNIPSVSSINRVLRLNLA SEKQQM . . . . . 137
fly/1-898      201 QEN VCTNDNIPSVSSINRVLRLNLA AQKEQ QSTGSGSSSTSAGNSISAKVS 250

human/1-422      138 . . . . .GA . . . . .DG 141
fly/1-898      251 VSIGGNVSNVASGSRGTLSSSTDLMQTATPLNSSSEG G A SNSGEGSEQEA 300

human/1-422      142 MYD KLRMLNG QTG . . . . . 154
fly/1-898      301 IYE KLRLLNT QHAAGPGPLEPARAAPLVGQSPNHLGTRSSHPQLVHGNHQ 350

human/1-422      155 . . . . .SWGTR . . . PGWYPGTSVPGQP TQ . . . . . 174
fly/1-898      351 ALQQHQQQSWPPRHYSGSWYP - TSLSEIPISSAPN IASVTAYASGPSLAH 399

human/1-422      175 . . . . .DGCCQQE . . . GGGEENTN 188
fly/1-898      400 SLSPNDIESLASIGHQRNCPVATEDIHLKKEL DGHQSD E TSGEGEENS N 449

human/1-422      189 S I S SNGEDSDEAQMRLQL KKRKLQRNRTSFTQEQ IEA LEKEFERTHYPDV 238
fly/1-898      450 GGA S N I G N T E D D Q A R L I L K R K L Q R N R T S F T N D Q I D S L E K E F E R T H Y P D V F 499

human/1-422      239 ARERLA AKID LPEARIQVWFSNRRAKWRREEKLRNQRRQASNT PSHIPIS 288
fly/1-898      500 ARERLAG KIG LPEARIQVWFSNRRAKWRREEKLRNQRRTPNST GASATSS 549

human/1-422      289 S S F S T S V Y Q P I P Q P T T P V S S F T S G S M L G R T D T A L T N T Y S A L P P M P S F T M A 338
fly/1-898      550 S T S A T A S L T D S P N S L S A C S S L L S G S A G G P S V S T I N G L S S . . . . . P S T L S T 594

human/1-422      339 N . N L P . . . . .MQPFVPSQTSSYS CMLPTSPSVNGRSYD . . . . .TYT 373
fly/1-898      595 NVNAPT LGAGIDSSSESTPIPHIRPSC . . . TSDNDNGRQSEDCRRVCSPC 641

human/1-422      374 PPHMQTHMNSQPMGTS GTTSTGLISP GVSVPVQVPGSEPDMSQYW PRLQ . 422
fly/1-898      642 PLGVGGH QNTHHIQSN GHAQGHALVPAIS . . . . . PRLNF 675

human/1-422      . . . . .
fly/1-898      676 NSGSFGAMYSNMHTALSMSDSYGAVTPIPSFNHSAVGPLAPPSP I PQQG 725

human/1-422      . . . . .
fly/1-898      726 DLTPSSLYPCHMTLRPPPMAPAHHH I VPGDGRPAGVGLGSGQSANLGAS 775

human/1-422      . . . . .
fly/1-898      776 CSGSGYEVLSAYALPPPPMASSSAADSSFS AASSASANVTPHHTIAQESC 825

human/1-422      . . . . .
fly/1-898      826 P S P C S S A S H F G V A H S S G F S S D P I S P A V S S Y A H M S Y N Y A S S A N T M T P S S A S 875

human/1-422      . . . . .
fly/1-898      876 G T S A H V A P G K Q Q F F A S C F Y S P W V 898

```

Alignment of an orthologous protein in *D.melanogaster* vs *H.sapiens*

Not suitable for global alignment

2 main regions of similarity

Better to use local alignment

Local alignment – Smith-Waterman algorithm

Example – align TCGA to GAC

0	-	T	C	G	A
-	0	0	0	0	0
G	0	-1	-1	1	-1
A	0	-1	-2	-1	2
C	0	-1	0	-2	0

Points for match = +1

Points for mismatch = -1

Points for a gap insertion = -2

Local alignment – Smith-Waterman algorithm

Example – align TCGA to GAC

0	-	T	C	G	A
-	0	0	0	0	0
G	0				
A	0				
C	0				

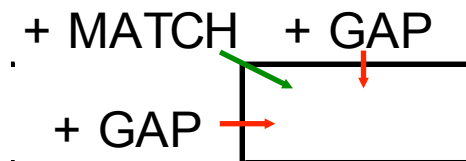
Points for match = +1

Points for mismatch = -1

Points for a gap insertion = -2

Local alignment – Smith-Waterman algorithm

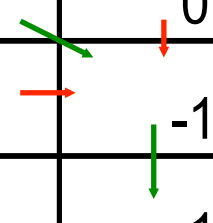
0	-	T	C	G	A
-	0	0	0	0	0
G	0	-1			
A	0				
C	0				



Points for match = +1
 Points for mismatch = -1
 Points for a gap insertion = -2

Local alignment – Smith-Waterman algorithm

0	-	T	C	G	A
-	0	0	0	0	0
G	0	-1			
A	0	-1			
C	0				



Points for match = +1

Points for mismatch = -1

Points for a gap insertion = -2

Local alignment – Smith-Waterman algorithm

0	-	T	C	G	A
-	0	0	0	0	0
G	0	-1			
A	0	-1			
C	0	-1			

Points for match = +1

Points for mismatch = -1

Points for a gap insertion = -2

Local alignment – Smith-Waterman algorithm

Example – align TCGA to GAC

0	-	T	C	G	A
-	0	0	0	0	0
G	0	-1	-1	1	-1
A	0	-1	-2	-1	2
C	0	-1	0	-2	0

Points for match = +1
Points for mismatch = -1
Points for a gap insertion = -2

Backtracking and final local alignment

0	-	T	C	G	A
-	0	0	0	0	0
G	0	-1	-1	1	-1
A	0	-1	-2	-1	2
C	0	-1	0	-2	0

GA
|
GA

Smith-Waterman – more details

<http://www.youtube.com/watch?v=IVRSFaGCGeE>

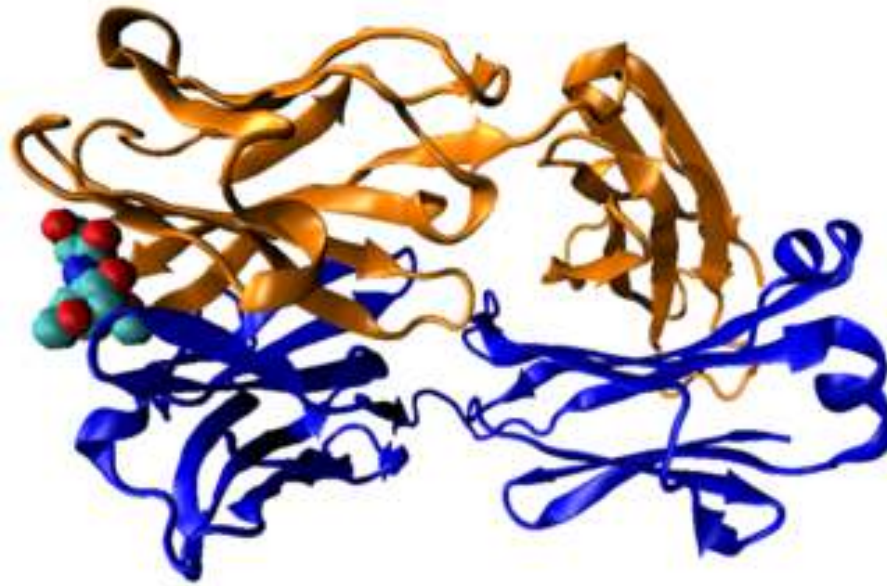
Dynamic programming

- Needleman-Wunsch and Smith-Waterman are a class of methods known as ‘Dynamic Programming’
- Guaranteed to give you the best possible alignment
- In biology, this algorithm is very inefficient because any 2 randomly selected DNA fragments in a database are unlikely to have any similarity
- Therefore, these methods take a long time to run

BLAST – Basic Local Alignment Search Tool

Background – BLAST

- Primarily designed to identify homologous sequences
 - Blast is a hashed seed-extend algorithm
 - Negative selection
 - Only some parts of a sequence are usually constrained



BLAST - Original version

Example:

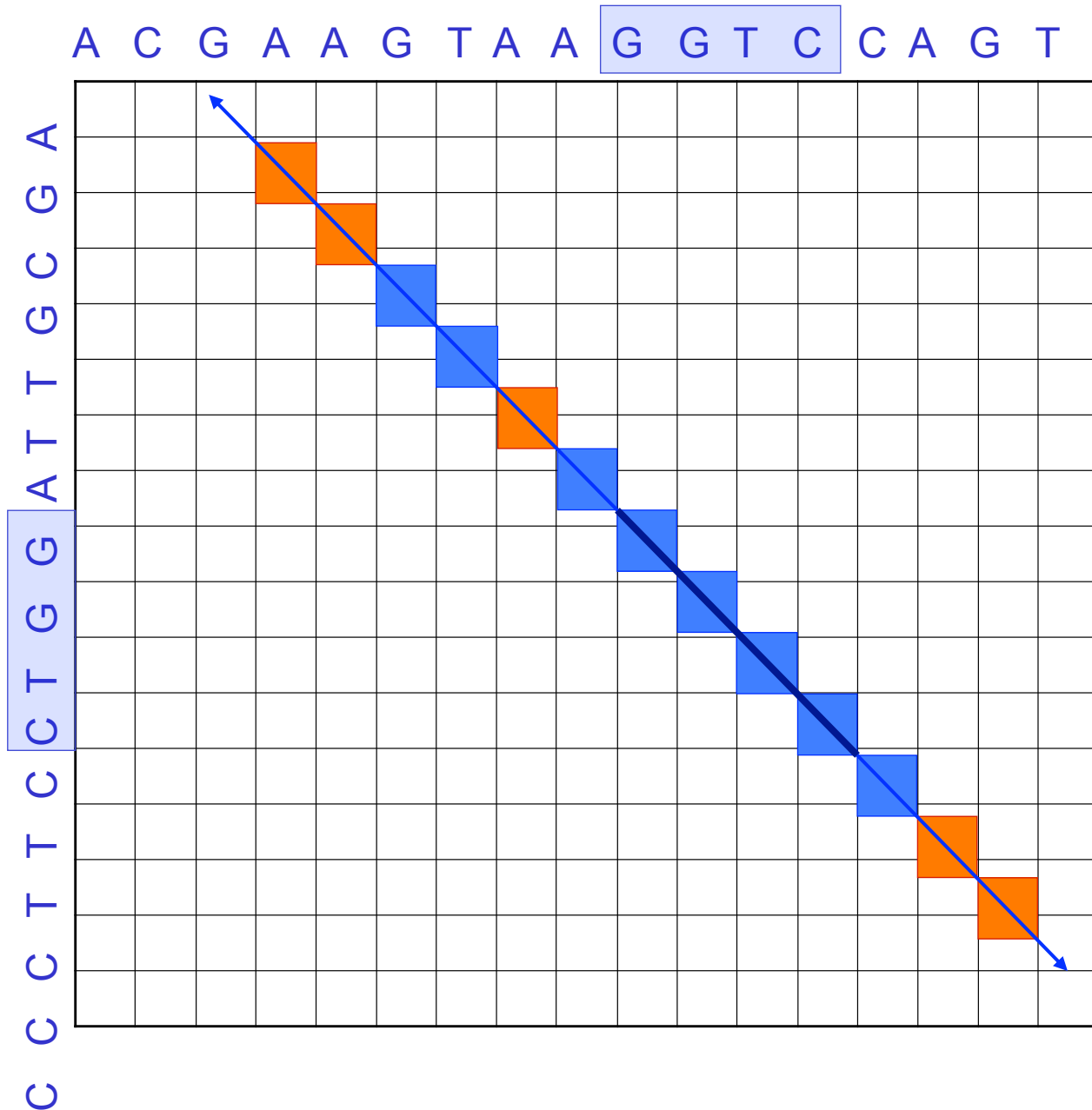
Seed size = 4,
No mismatches in seed

The matching word GGTC
initiates an alignment

Extension to the left and right
with no gaps until alignment
score falls below 50%

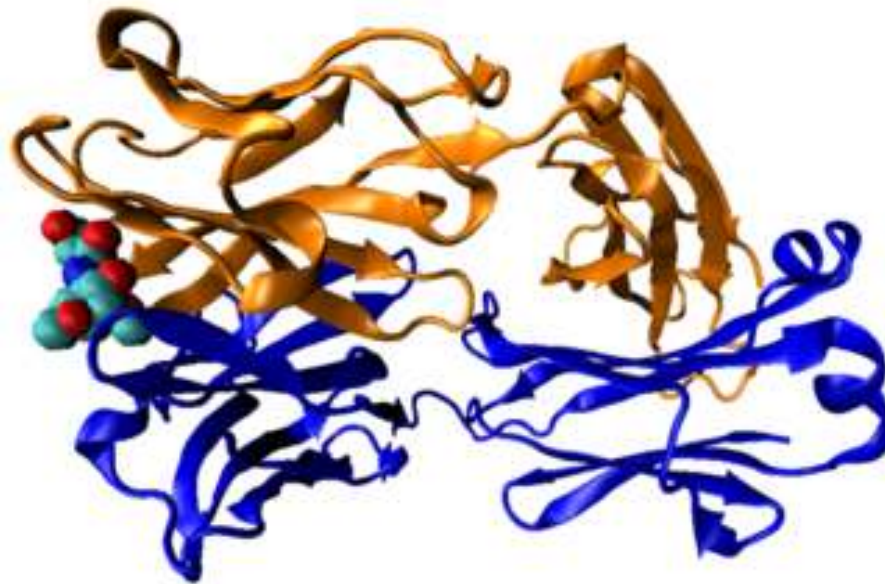
Output:

GTAAGGTCC
GTTAGGTCC



BLAST - Original algorithm

- Finding seeds significantly increases the speed of BLAST compared to doing a full local alignment over a whole sequence
- Will not guarantee the best solution
- BLAST first finds highly conserved or identical sequences which are then extended with a local alignment.



BLAST – Speed (or lack thereof)

- Typically BLAST will take approximately 0.1 – 1 second to search 1 sequence against a database
- Depends on size of database, e-value cutoff and number of hits to report selected
- 60 million reads equates to 70 CPU days!
- Even on multi-core systems this is too long!
- Especially if you have multiple samples!
- This is still true of FPGA and SIMD (vectorised) implementations of BLAST

When **NOT** to use *BLAST*

- A typical situation: you have lots DNA sequences and want to extend it or find where on a genome it maps.
- In other words, you want an **exact** or **near-exact** match to a sequence that is part of an **assembled genome**.
- Short reads require very fast algorithms for finding near-exact matches in genomic sequences:
 - BLAT
 - Highly recommended: the BLAT paper (Kent WJ (2003) *Genome Res* 12:656-64) – very well written
 - SOAP
 - Bowtie/Bowtie 2
 - MAQ
 - BWA
 - Shrimp2

Contents

- **Alignment algorithms for short-reads**
 - Background – Blast (why can't we use it?)
 - **Adapting hashed seed-extend algorithms to work with shorter reads**
 - Indel detection
 - Suffix/Prefix Tries
 - Other alignment considerations
 - Typical alignment pipeline
 - New methods of variant calling

Adapting hashed seed-extend algorithms to work with shorter reads

- Improve seed matching sensitivity
 - Allow mismatches within seed
 - BLAST
 - Allow mismatches + Adopt spaced-seed approach
 - ELAND, SOAP, MAQ, RMAP, ZOOM
 - Allow mismatches + Spaced-seeds + Multi-seeds
 - SSAHA2, BLAT, ELAND2
- Above and/or Improve speed of local alignment for seed extension
 - Single Instruction Multiple Data
 - Shrimp2, CLCBio
 - Reduce search space to region around seed

Hashed seed-extend algorithms

- These are most similar to BLAST
- Are not designed to work with large databases
- **2 step process**
 - Identify a match to the seed sequence in the reference
 - Extend match using sensitive (but slow) Smith-Waterman algorithm (dynamic programming)

Seed-extend algorithm

Reference sequence:

...ACTGGGTCATCGTACGATCGATCGATCGATCGATCGGCTAGCTAGCTA...

Short read:

GTCATCGTACGATCGATAGATCGATCGGCTA

Note that the short read has 1 difference wrt to reference

Seed-extend algorithm

Reference sequence:

...ACTGGGTCATCGTACGATCGATCGATCGATCGATCGGCTAGCTAGCTA...

Short read:

GTCATCGTACG ATCGATAGATCG ATCGATCGGCTA

11bp word

11bp word

11bp word

The algorithm will try to match each word to the reference. If there is a match at with any single word it will perform a local alignment to extend the match

Seed-extend algorithm

Reference sequence:

Seed Extend with Smith Waterman
...ACTGGG **GTCATCGTACG** **ATCGATCGATCGATCGATCGGCTA** GCTAGCTA...
 GTCATCGTACG **ATCGA** **ACGATCGATCGATCGGCTA**

Short read:

GTCATCGTACG **ATCGATAGATCG** **ATCGATCGGCTA**

Here the algorithm is able to match the short read with a word length of 11bp

Seed-extend algorithm

Reference sequence:

...ACTGGGTCATCGTACGATCGATCGATCGATCGATCGGCTAGCTAGCTA...

Short read:

GTCATCGTACGATCGATCGATCGATCGGCAA

Note that the short read has 3 differences
Possibly sequencing errors, possibly SNPs

Seed-extend algorithm

Reference sequence:

...ACTGGGTCATCGTACGATCGATCGATCGATCGATCGGCTAGCTAGCTA...

Short read:

GTCATCGTACG ATCGATCGATCG
11bp word 11bp word 11bp word

Note that the short read has 3 differences

Seed-extend algorithm

Reference sequence:

...ACTGGGTCATCGTACGATCGATCGATCGATCGATCGGCTAGCTAGCTA...

Short read:

GTCATCGTACG ATCGATCGATCG
ATCGATCGGCA

No seeds match

Therefore the algorithm would find no hits at all!

–
📄
✕

Downloaded from <http://ajphaphysocpharm.sagepub.com/> at 11:00 11 November 2014



Adapting hashed seed-extend algorithms to work with shorter reads

- Improve seed matching sensitivity
 - **Allow mismatches within seed**
 - **BLAST**
 - Allow mismatches + Adopt spaced-seed approach
 - ELAND, SOAP, MAQ, RMAP, ZOOM
 - Allow mismatches + Spaced-seeds + Multi-seeds
 - SSAHA2, BLAT, ELAND2
- Above and/or Improve speed of local alignment for seed extension
 - Single Instruction Multiple Data
 - Shrimp2, CLCBio
 - Reduce search space to region around seed

Adapting hashed seed-extend algorithms to work with shorter reads

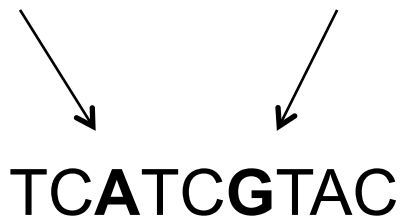
- Improve seed matching sensitivity
 - **Allow mismatches within seed**
 - **BLAST**
 - **Allow mismatches + Adopt spaced-seed approach**
 - **ELAND, MAQ, RMAP, ZOOM**
 - Allow mismatches + Spaced-seeds + Multi-seeds
 - SSAHA2, BLAT, ELAND2
- Above and/or Improve speed of local alignment for seed extension
 - Single Instruction Multiple Data
 - Shrimp2, CLCBio
 - Reduce search space to region around seed

Consecutive seed

Consecutive seed 9bp allowing no mismatches:

ACTCCCATCGTCATCGTACTAGGGATCGTAACA Reference sequence

CCACTGTCCTCTACATAGGAACGA SNP 'heavy' read



TCATCGTAC

TCCTCTAC Cannot find seed match due to A->C SNP
and G->C SNP

Even allowing for 2 mismatches in
the seed - no seeds match.

No hits!

Spaced seeds

To increase sensitivity we can use spaced-seeds:

111111111

Consecutive seed template with *length* 9bp

ACTATCATCGTACACAT

Reference

TCATCGTAC

Query

11001100110011001

Spaced-seed template with *weight* 9bp

ACTATCATCGTACACAT

Reference

ACTCTCA^CCGTACACAT

Query

Spaced seeds

Spaced seed with weight 9bp and no mismatches:

ACTCCCATTTGTCATCGTACTTGGGATCGTAACA Reference sequence

CCACTGTATCGTACATGGGAACGA SNP 'heavy' read



CCATTGTCATCGTACAT

CCXXTGXXATXXTAXXT

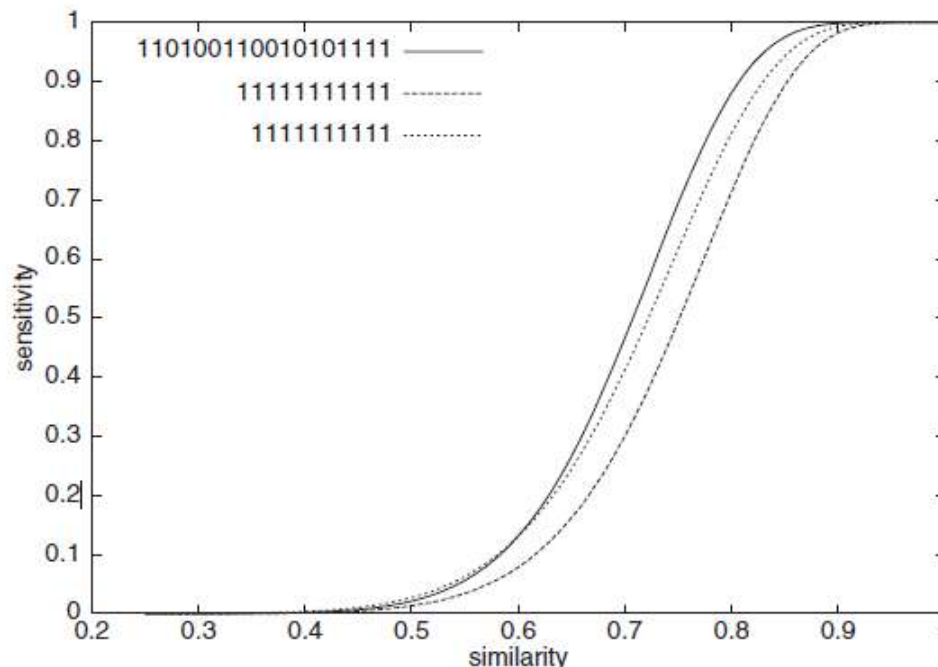
Despite SNPs – seed matched with 0 mismatches

Can now extend with Smith-Waterman or other local alignment

Spaced seeds

Spaced seeds:

- A seed template '111010010100110111' is 55% more sensitive than BLAST's default template '1111111111' for two sequences of 70% similarity
- Typically seeds of length ~30bp and allow up to 2 mismatches in short read datasets



Contents

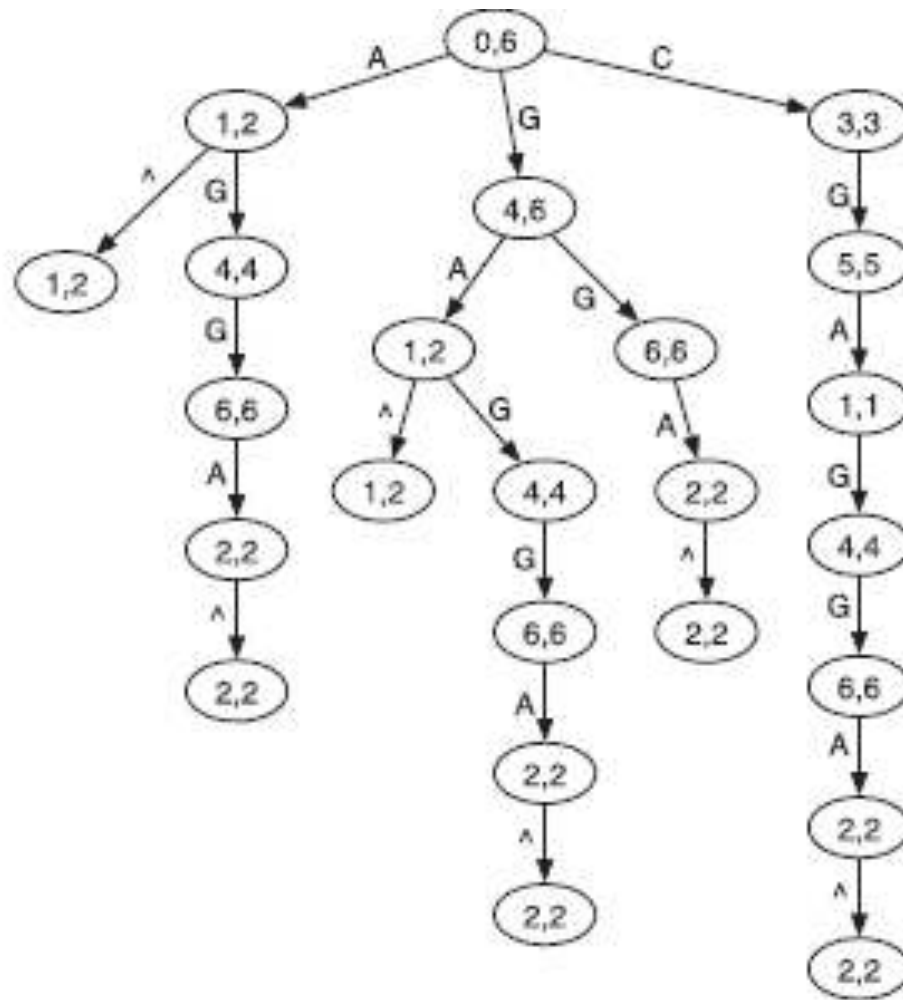
- **Alignment algorithms for short-reads**
 - Background – Blast (why can't we use it?)
 - Adapting hashed seed-extend algorithms to work with shorter reads
 - **Suffix/Prefix Tries**
 - Indel detection
 - Other alignment considerations
 - Typical alignment pipeline
 - New methods of SNP calling

Suffix-Prefix Trie

- Trie – data structure which stores the suffixes (i.e. ends of a sequence)
- A family of methods which uses a Trie structure to search a reference sequence
 - Bowtie
 - BWA aln (<70bp reads) and MEM algorithm (>70bp reads)
 - SOAP version 2
- Key advantages:
 - Alignment of multiple copies of an identical sequence in the reference only needs to be done once
 - Use of an FM-Index to store Trie can drastically reduce memory requirements (e.g. Human genome can be stored in 2Gb of RAM)
 - Burrows Wheeler Transform to perform fast lookups

Suffix Trie

Read
AGGAGC



Heng Li & Nils Homer.
Sequence alignment
algorithms for next-
generation sequencing.
Briefings in
Bioinformatics. Vol 11.
No 5. 473 483, 2010

Burrows-Wheeler Algorithm

- Encodes data so that it is easier to compress
- Burrows-Wheeler transform of the word BANANA
- Can later be reversed to recover the original word

Transformation				
Input	All Rotations	Sorting All Rows in Alphabetical Order by their first letters	Taking Last Column	Output Last Column
^BANANA	^BANANA ^BANANA A ^BANAN NA ^BANA ANA ^BAN NANA ^BA ANANA ^B BANANA ^	ANANA ^B ANA ^BAN A ^BANAN BANANA ^ NANA ^BA NA ^BANA ^BANANA ^BANANA	ANANA ^B ANA ^BAN A ^BANAN BANANA ^ NANA ^BA NA ^BANA ^BANANA ^BANANA	BNN^AA A

More Burrows-Wheeler

Input	SIX.MIXED.PIXIES.SIFT.SIXTY.PIXIE.DUST.BOXES
Burrows-Wheeler Output	TEXYDST.E.IXIXIXXSSMPPS.B..E.S.EUSFXDIIIOIIT

Repeated characters mean that it is easier to compress

Suffix Trie for a bacterial genome would be $> 1\text{Tb}$

We have to compress it

Use FM-Index/BW transform to do this compression

Bowtie/BWA example

Reference



BWT(Reference)

Query:

AATGATACGGCGACCACCGAGATCTA

Bowtie/BWA example

Reference



BWT(Reference)



Query:

AATGATACGGCGACCACCGAGATCTA



Bowtie/BWA example

Reference



BWT(Reference)

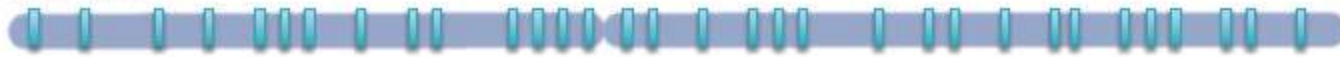


Query:

AATGATACGGCGACCACCGAGATCTA

Bowtie/BWA example

Reference



BWT(Reference)



Query:

AATGATACGGCGACCACCGAGATCTA

Bowtie/BWA example

Reference



BWT(Reference)



Query:

AATGATACGGCGAC **CACCGAGATCTA**

Bowtie/BWA example

Reference



BWT(Reference)



Query:

AATGATACGGCGACCAACGAGATCTA

Bowtie/BWA example

Reference



BWT(Reference)

Query:

AATGATACGGCGACCAACGAGATCTA

Bowtie/BWA example

Reference



BWT(Reference)



Query:

AATG TACGGCGACCAACGAGATCTA

Bowtie/BWA example

Reference



BWT(Reference)



Query:

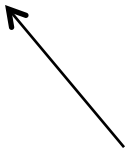
AATGTTACGGCGACCACCGAGATCTA

Bowtie/Soap2 vs. BWA

- Bowtie 1 and Soap2 cannot handle gapped alignments
 - No indel detection => Many false SNP calls

Bowtie/Soap2:

ACTCCCATTTGTCATCGTACTTGGGATC**G**TAACA Reference
CCATTGTCATCGTACTTGGGATC**TA**
TCATCGTACTTGGGATC**TA**
TTGGGATC**TA**



False SNPs

N.B. Bowtie2 can handle gapped alignments

Bowtie/Soap2 vs. BWA

- Bowtie 1 and Soap2 cannot handle gapped alignments
 - No indel detection => Many false SNP calls

BWA:

```
ACTCCCATTTGTCATCGTACTTTGGGATCGTAACA Reference
      CCATTGTCATCGTACTTTGGGATC-TA
        TCATCGTACTTTGGGATC-TA
          TTGGGATC-TA
```

N.B. Bowtie2 can handle gapped alignments

Comparison

Hash referenced spaced seeds

- Requires ~50Gb of memory
- Runs 30-fold slower
- Is much simpler to program
- Most sensitive

Indexed Suffix/Prefix Trie

- Requires <2Gb of memory
- Runs 30-fold faster
- Is much more complicated to program
- Least sensitive

There are limits however

With longer 100-300 bp reads, multiple indels or variable regions longer than a few bp are likely to be missed

ACTCCCATTGTCATCGTACTTGGGATC**G**TAACA Reference
CCATTGTCA**ACCATCTAGTAGCT**-TA
TCA**ACCATCTAGTAGCT**-TA
ACCATCTA-TA

You only find what you are looking for

- What happens if there are SNPs and Indels in the same region?

Let's assume that the SNP caller made this call of a single SNP:

ATGTA**TGTA**
ATGTG**TGTA**

and the indel caller produced this call of a 3 base deletion:

ATGTATGTA
ATGT---TA

Should we assume this is a heterozygous SNP opposite a heterozygous Indel or a more complex locus?

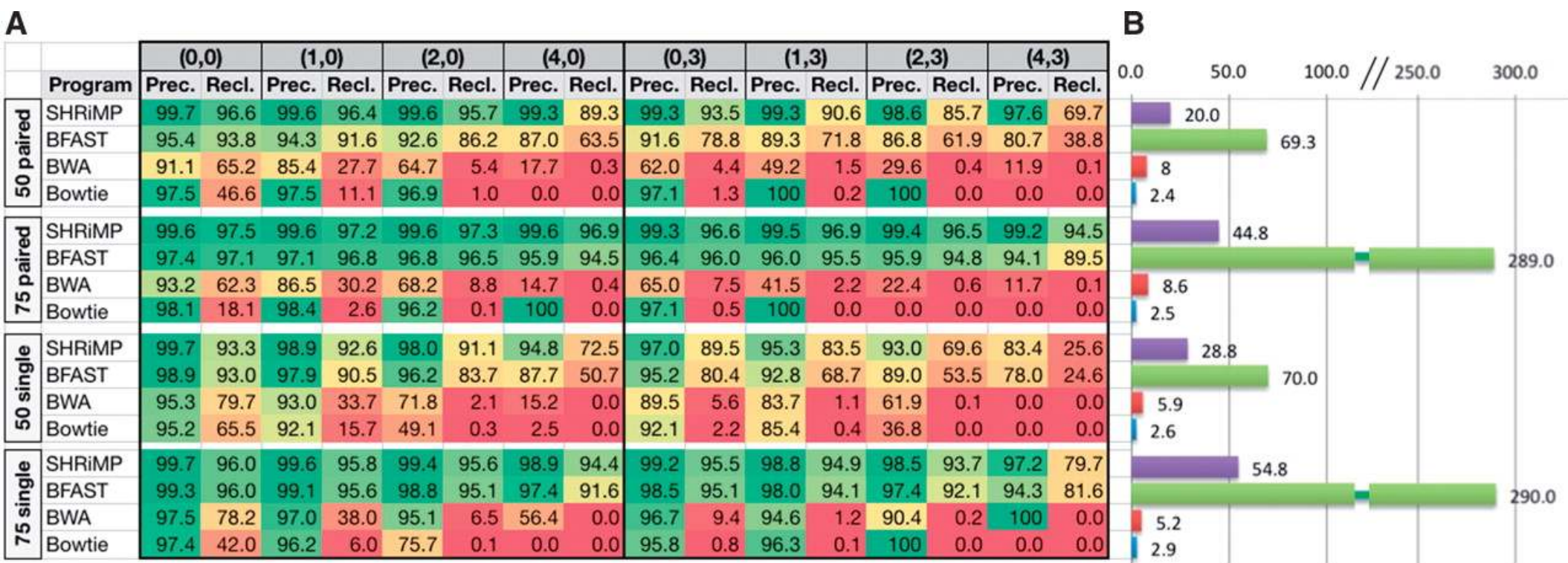
Comparison

- Bowtie's reported 30-fold speed increase over hash-based methods with small loss in sensitivity
- Limitations to Trie-based approaches:
 - Only able to find alignments within a certain 'edit distance'
 - Important to quality clip reads (-q in BWA)
 - Non-A/C/G/T bases on reads are often treated as mismatches
 - Make sure Ns are removed!

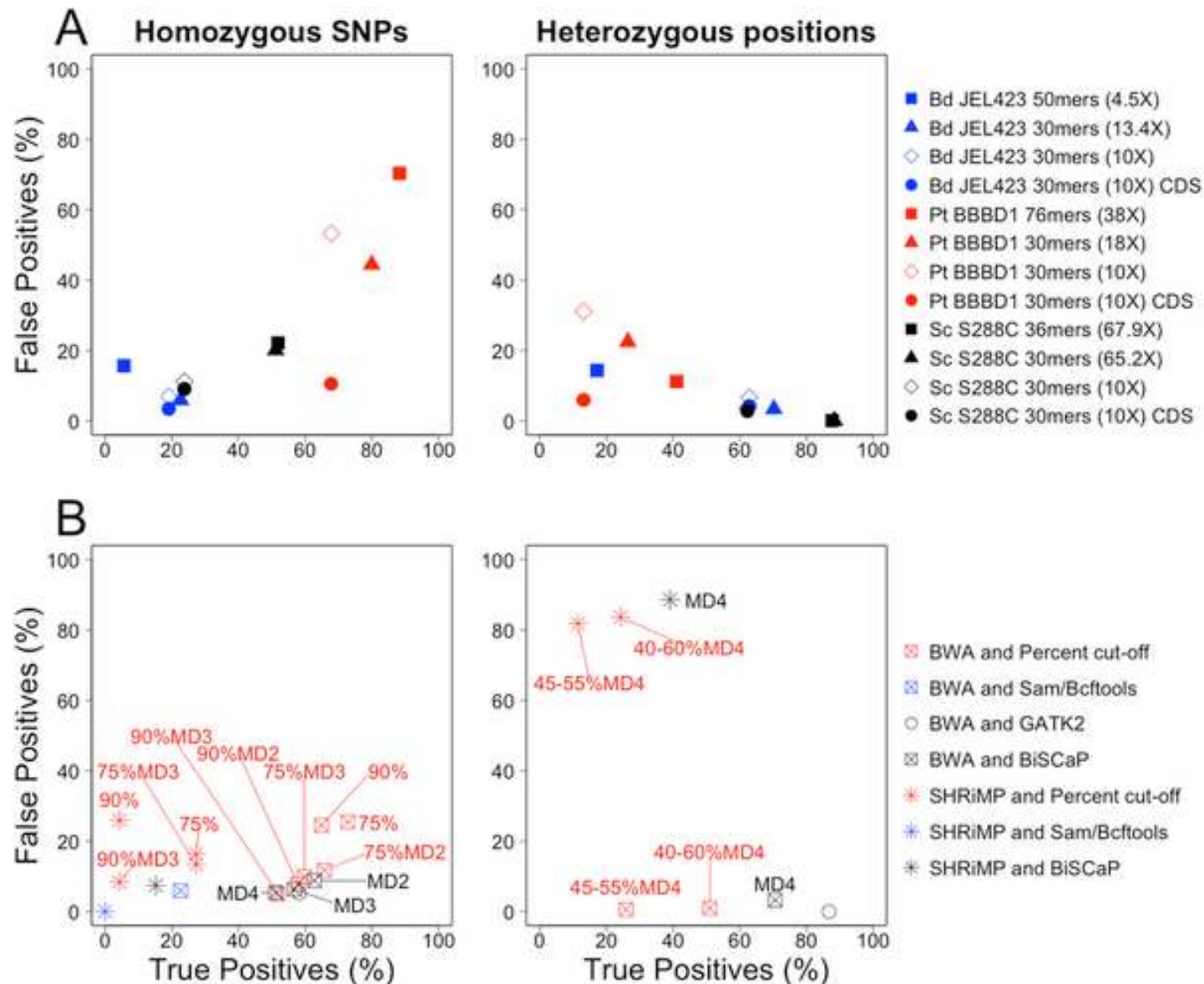
Hash based approaches are more suitable for divergent alignments

- Rule of thumb:
 - <2% divergence -> Trie-based
 - E.g. human alignments
 - >2% divergence -> Seed-extend based approach
 - E.g. wild mouse strain alignments

Precision and recall by amount of variation for 4 datasets, by polymorphism: (number of SNPs, Indel size)



False discovery rates for variants were ascertained using cFDR for three fungal NGS datasets



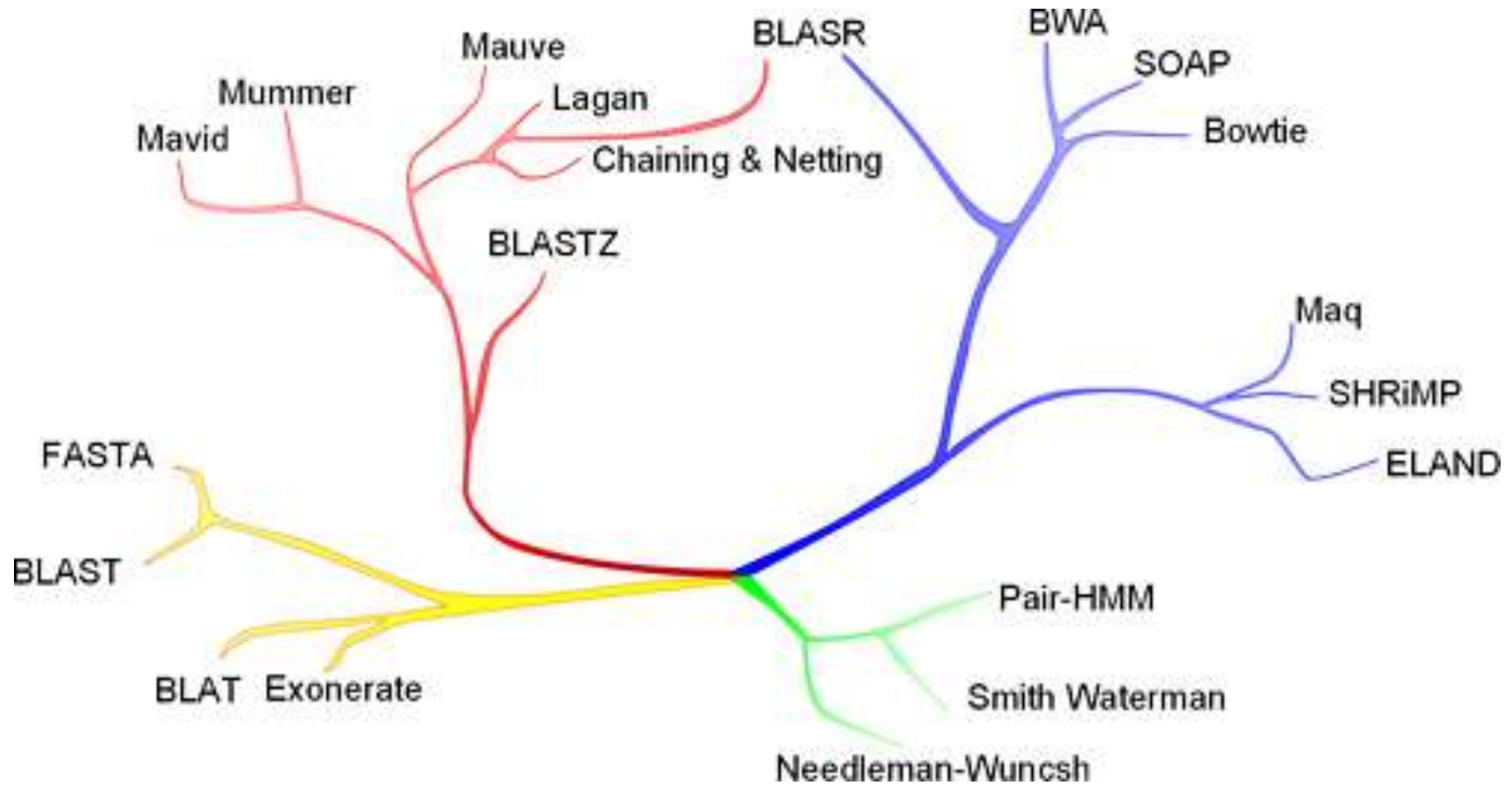
Summary of open-source short read alignment programs

Program	Algorithm	SoLID	Long reads	Gapped alignment	Paired-end	Quality scores used?
Bfast	Hashing ref	Yes	No	Yes	Yes	No
Bowtie2*	FM-Index	Yes	Yes	Yes	Yes	Yes
Blat	Hashing ref	No	Yes	Yes	No	No
BWA	FM-Index	Yes	Yes	Yes	Yes	No
MAQ	Hashing reads	Yes	No	Yes	Yes	Yes
Mosaik	Hashing ref	Yes	Yes	Yes	Yes	No
Novoalign	Hashing ref	No	No	Yes	Yes	Yes
Shrimp2	Hashing ref	Yes	Yes	Yes	Yes	Yes
SOAP2	FM-Index	No	No	No	Yes	Yes
SSAHA2	Hashing ref.	No	No	No	Yes	Yes

Heng Li & Nils Homer. Sequence alignment algorithms for next-generation sequencing. Briefings in Bioinformatics. Vol 11. No 5. 473 483, 2010

* Bowtie1 does not support gapped alignments

Aligner phylogeny



Whole genome

Pairwise heuristic

Large data set aligners

Sensitive global aligners

Sequence read aligners I use

- Genomic alignments BWA-Mem
 - Scales well with read lengths and will tolerate more errors as read lengths increase
- Bowtie2/Tophat for RNA-seq alignment
 - Splice aware and fits into a nice eco-system of tools to perform abundance expression (Cufflinks) and visualisation (cummeRbund)
- BLASR
 - Designed for PacBio reads

Alignment format for short reads – Sequence AlignMent (SAM format)

- Plain text format – human readable (sort-of)
- Eleven mandatory fields and a variable amount of optional fields.
- The optional fields are a key-value pair of TAG:TYPE:VALUE. These store extra information
- Can be converted to Binary AlignMent format (BAM) to save space and speed up look-up operations using SAMTools

Alignment format for short reads – Sequence Alignment (SAM format)

Table 1. Mandatory fields in the SAM format

No.	Name	Description
1	QNAME	Query NAME of the read or the read pair
2	FLAG	Bitwise FLAG (pairing, strand, mate strand, etc.)
3	RNAME	Reference sequence NAME
4	POS	1-Based leftmost POSition of clipped alignment
5	MAPQ	MAPping Quality (Phred-scaled)
6	CIGAR	Extended CIGAR string (operations: MIDNSHP)
7	MRNM	Mate Reference NaMe ('=' if same as RNAME)
8	MPOS	1-Based leftmost Mate POSition
9	ISIZE	Inferred Insert SIZE
10	SEQ	Query SEQUENCE on the same strand as the reference
11	QUAL	Query QUALity (ASCII-33=Phred base quality)

SAM format – Optional fields

Tag	Type	Description
X?	?	Reserved fields for end users (together with Y? and Z?)
AM	i	The smallest template-independent mapping quality of fragments in the rest
AS	i	Alignment score generated by aligner
BQ	Z	Offset to base alignment quality (BAQ), of the same length as the read sequence. At the i -th read base, $BAQ_i = Q_i - (BQ_i - 64)$ where Q_i is the i -th base quality.
CM	i	Edit distance between the color sequence and the color reference (see also NM)
CQ	Z	Color read quality on the original strand of the read. Same encoding as QUAL; same length as CS.
CS	Z	Color read sequence on the original strand of the read. The primer base must be included.
E2	Z	The 2nd most likely base calls. Same encoding and same length as QUAL.
FI	i	The index of fragment in the template.
FS	Z	Fragment suffix.
LB	Z	Library. Value to be consistent with the header RG-LB tag if @RG is present.
H0	i	Number of perfect hits
H1	i	Number of 1-difference hits (see also NM)
H2	i	Number of 2-difference hits
HI	i	Query hit index, indicating the alignment record is the i -th one stored in SAM
IH	i	Number of stored alignments in SAM that contains the query in the current record
MD	Z	String for mismatching positions. <i>Regex</i> : $[0-9]+((([ACGTN] \^[ACGTN])+)[0-9]+)^1$
MQ	i	Mapping quality of the mate/next fragment
NH	i	Number of reported alignments that contains the query in the current record
NM	i	Edit distance to the reference, including ambiguous bases but excluding clipping
OQ	Z	Original base quality (usually before recalibration). Same encoding as QUAL.
OP	i	Original mapping position (usually before realignment)
OC	Z	Original CIGAR (usually before realignment)
PG	Z	Program. Value matches the header PG-ID tag if @PG is present.
PQ	i	Phred likelihood of the template, conditional on both the mapping being correct
PU	Z	Platform unit. Value to be consistent with the header RG-PU tag if @RG is present.
Q2	Z	Phred quality of the mate/next fragment. Same encoding as QUAL.
R2	Z	Sequence of the mate/next fragment in the template.
RG	Z	Read group. Value matches the header RG-ID tag if @RG is present in the header.
SM	i	Template-independent mapping quality
TC	i	The number of fragments in the template.
U2	Z	Phred probility of the 2nd call being wrong conditional on the best being wrong. The same encoding as QUAL.
UQ	i	Phred likelihood of the fragment, conditional on the mapping being correct

SAM output

```
amaxwell@pinfish:~/ecoli/illum/bwamem/i4m1_1/seqAssist1/result
[amaxwell@pinfish result]$ head -10 aln.sam
@SQ      SN:gi|49175990|ref|NC_000913.2| LN:4639675
HWI-EAS397:8:1:1067:18713#CTTGTA      16      gi|49175990|ref|NC_000913.2|      2909788 35      13S36M      *      0      0      TACAATAACCCCCCGCCCCTGGGTAATAAGCCGACAATCTCATCTCCA      BBBB BBBB BBBB BBBB^aa\^^\V_dkwdfad
affccffcY|adcL^Y      NM:i:1 AS:i:31 XS:i:0
HWI-EAS397:8:1:1070:11813#CTTGTA      16      gi|49175990|ref|NC_000913.2|      4516890 24      49M      *      0      0      CGTCGGTGCTAAAGCAGGTACGCGCTGGCTGTTTTACGCGTATGACAGG      BaYa[\Kaa0GVHLZLIXH`]VaJR]V]V_J]
Qca|\aaS[J_aacccc      NM:i:2 AS:i:39 XS:i:30
HWI-EAS397:8:1:1070:11813#CTTGTA      256      gi|49175990|ref|NC_000913.2|      278711 0      30M19H      *      0      0      *      *      NM:i:0 AS:i:30
HWI-EAS397:8:1:1070:11813#CTTGTA      256      gi|49175990|ref|NC_000913.2|      290182 0      30M19H      *      0      0      *      *      NM:i:0 AS:i:30
HWI-EAS397:8:1:1070:11813#CTTGTA      272      gi|49175990|ref|NC_000913.2|      1049415 0      19H30M      *      0      0      *      *      NM:i:0 AS:i:30
HWI-EAS397:8:1:1121:19907#CTTGTA      0      gi|49175990|ref|NC_000913.2|      953460 60      49M      *      0      0      AGAGAGAAGCAAATGCCGCCAACCAGTTTTGCCATGCCGAAGGGCATTG      SIaJaL^ZT`[da^WcacfK^acSacffffff[
cY^dcdd^f]\|cffff      NM:i:0 AS:i:49 XS:i:0
HWI-EAS397:8:1:1138:20219#CTTGTA      4      *      0      0      *      *      0      0      TAAAAAAAGGCCATAACCTGCCGCTTTGTATAATAAAAAAGGGCCCGCG      Yd[cK^__wdc\ddada]df[fccff_f_ffffcfcdBBBBBBBBBBB      A
S:i:0 XS:i:0
HWI-EAS397:8:1:1149:21173#GTTGTA      4      *      0      0      *      *      0      0      GTTCGTTTTTTTTGTACAGTGCCAACATCTGCTCCCGCCAATGGTGCG      Scc_ZccccaVZR^[W_K_M_\^VSQ]\\\aS^aVaXL]WY^^IaaaB      A
S:i:0 XS:i:0
HWI-EAS397:8:1:1152:20684#CTTGTA      4      *      0      0      *      *      0      0      GCTGATGCGTACGTGAAAGGGCGGGACACCTTGGATGTATGGTTTGAGT      aQ^aZ^M0aX0X_V\_cc_BBBBBBBBBBBBBBBBBBBBBBBBBBBBBBBB      A
S:i:0 XS:i:0
[amaxwell@pinfish result]$
```


Contents

- **Alignment algorithms for short-reads**
 - Background – Blast (why can't we use it?)
 - Adapting hashed seed-extend algorithms to work with shorter reads
 - Indel detection
 - Suffix/Prefix Tries
 - **Other alignment considerations**
 - Typical alignment pipeline
 - New methods of SNP calling

Other alignment considerations

- Indel detection
- Effect of paired-end alignments
- Using base quality to inform alignments
- PCR duplicates
- Methylation experiments – bisulfite treated reads
- Multi-mapping reads
- Aligning spliced-reads from RNA-seq experiments
- Local realignment to improve SNP/Indel detection
- Platform specific errors
- Unmapped reads

Indel detection

Spaced seed with weight 9bp and no mismatches:

ACTCCCATTGTCATCGTACTTGGGATCGTAACA Reference sequence

CCATTGTCATGTACTTGGGATCGT

Read containing a
deletion



CCATTGTCATCGTACAT

CCXXTGXXATXXACXXG Seed not matched due to frame shift caused
by gap

No seed match. No alignment!

Indel detection

Reference sequence:

Seed Extend with Smith Waterman

...ACTGGG **GTCATCGTACG** ATCGATCGATCGATCGATCGGCTA GCTAGCTA...

GTCATCGTACG **ATCGA-CGATCGATCGATCGGCTA**

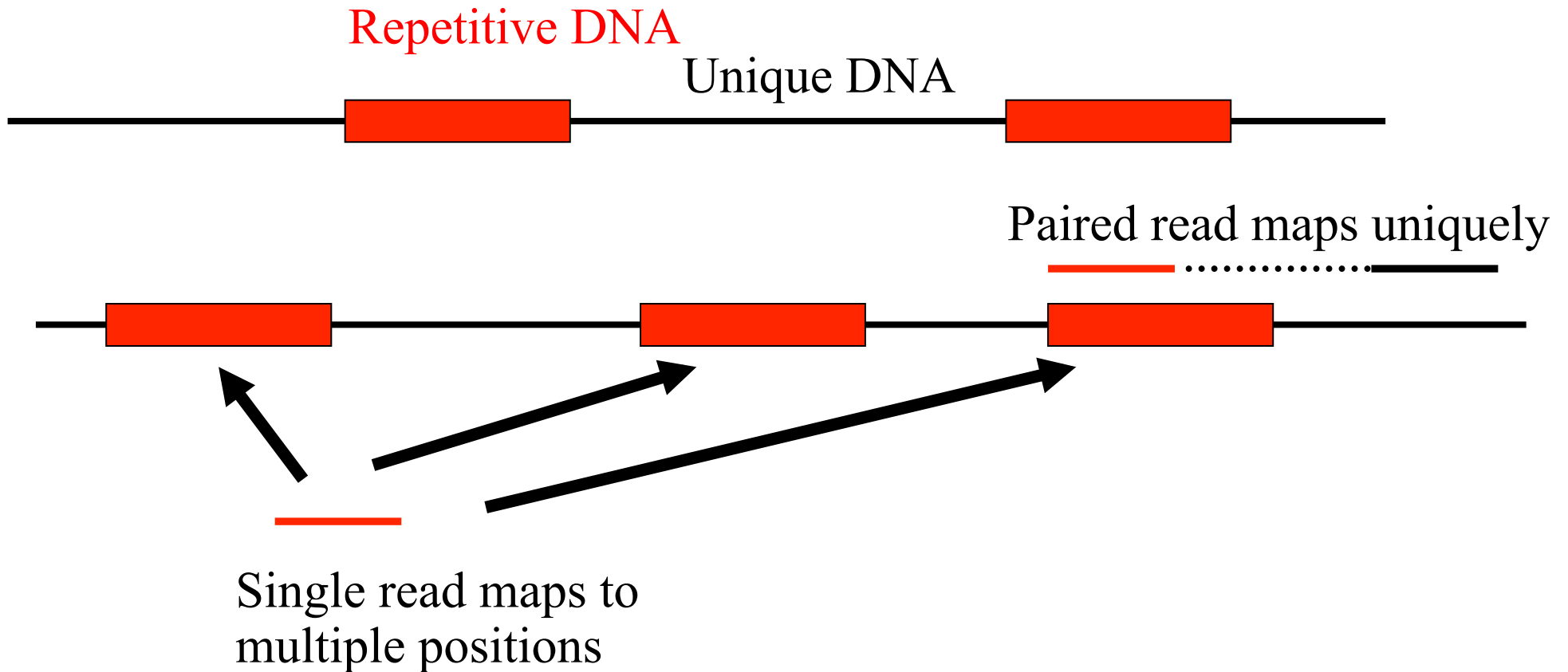
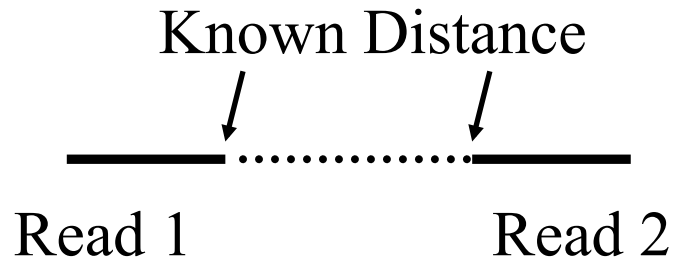
Most alignment programs can only detect gaps in
Smith-Waterman phase
once a seed has been identified. Some algorithms (e.g.
Bowtie) do not allow gaps at this stage to improve
speed

**This reduces sensitivity especially with multiple
insertions in a small region**

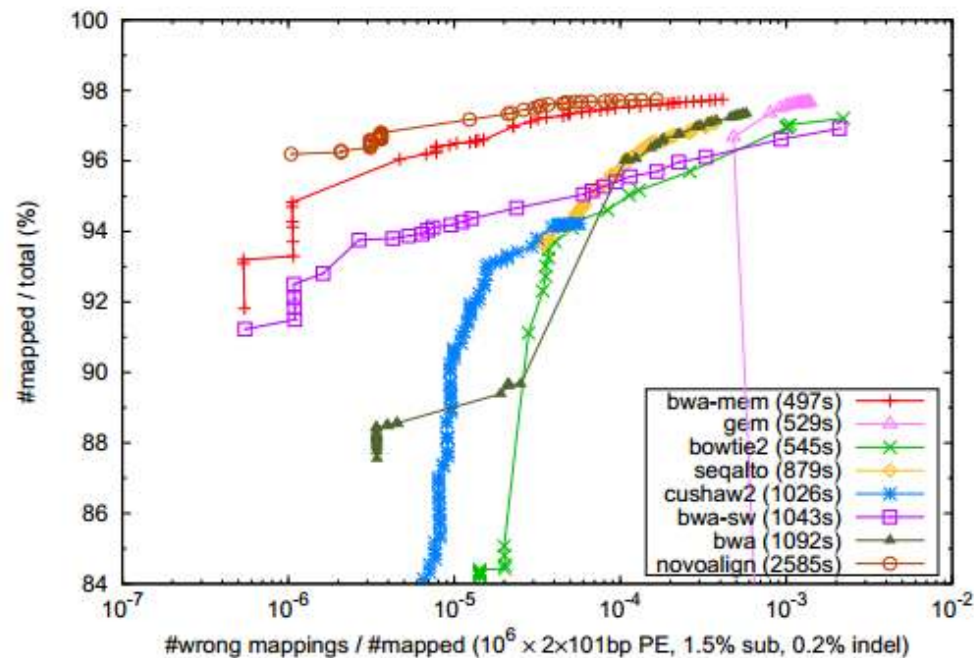
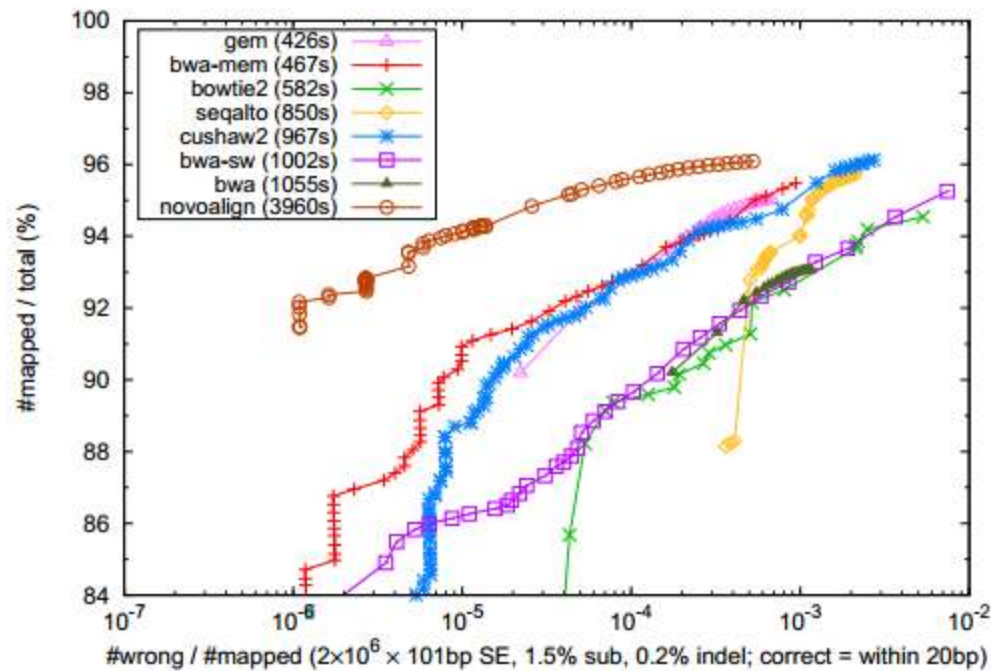
Indel detection

- Some algorithms do allow gaps within seed
 - Indel seeds for homology search *Bioinformatics* (2006) 22(14): e341-e349
doi:10.1093/bioinformatics/btl263
 - Weese D, Emde AK, Rausch T, et al. RazerS—fast read mapping with sensitivity control. *Genome Res* 2009;19:1646–54
 - Rumble SM, Lacroute P, Dalca AV, et al. SHRiMP: accurate mapping of short color-space reads. *PLoS Comput Biol* 2009;5:e1000386
- Use of multiple seeds
 - Especially useful for longer reads (>50bp)
 - Li R, Li Y, Kristiansen K, et al. SOAP: short oligonucleotide alignment program. *Bioinformatics* 2008;24:713–4
 - Jiang H, Wong WH. SeqMap: mapping massive amount of oligonucleotides to the genome. *Bioinformatics* 2008;24: 2395–6

Paired-end reads are important



Effect of paired-end alignments

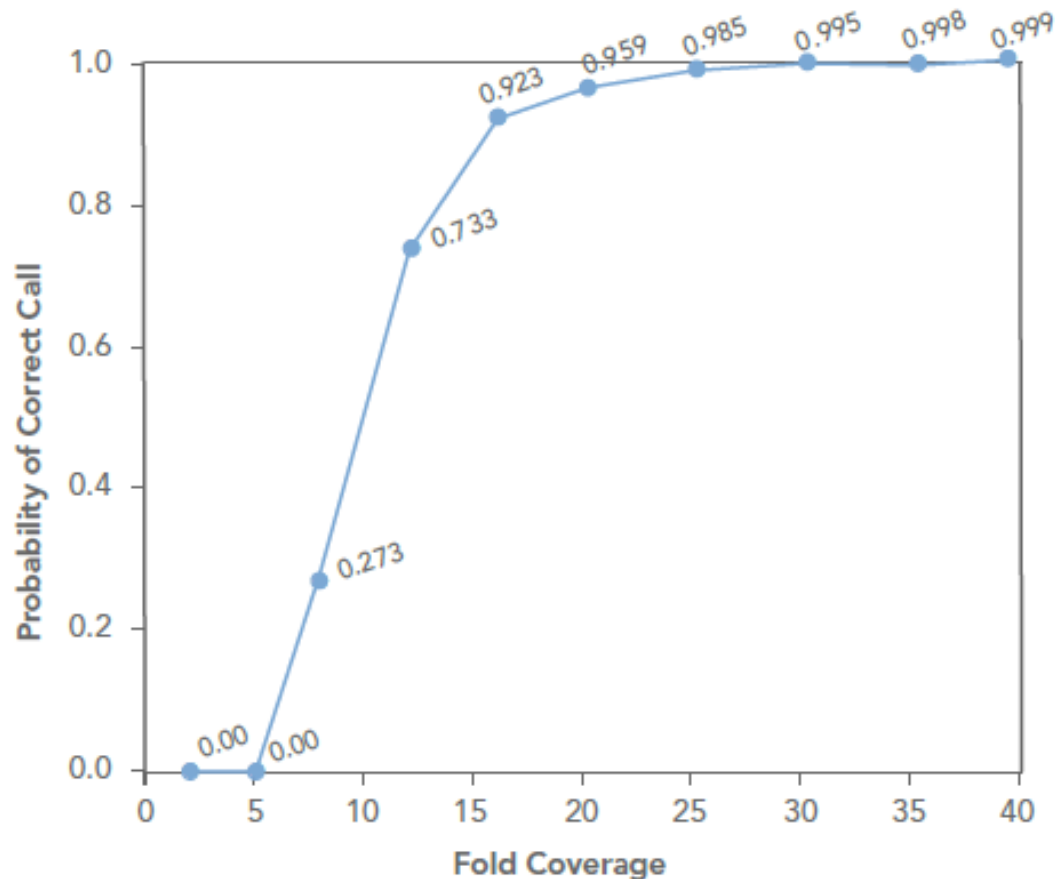


BWA-MEM

<http://arxiv.org/pdf/1303.3997v2.pdf>

Effect of coverage on SNP call accuracy

- Depends on ploidy
- Bacterial genomes can get away with 10-20x
- For human genomes and other diploids 30x
- Poly-ploids (e.g wheat) may need much higher coverage



*Source – Illumina Tech Note
Human diploid sample*

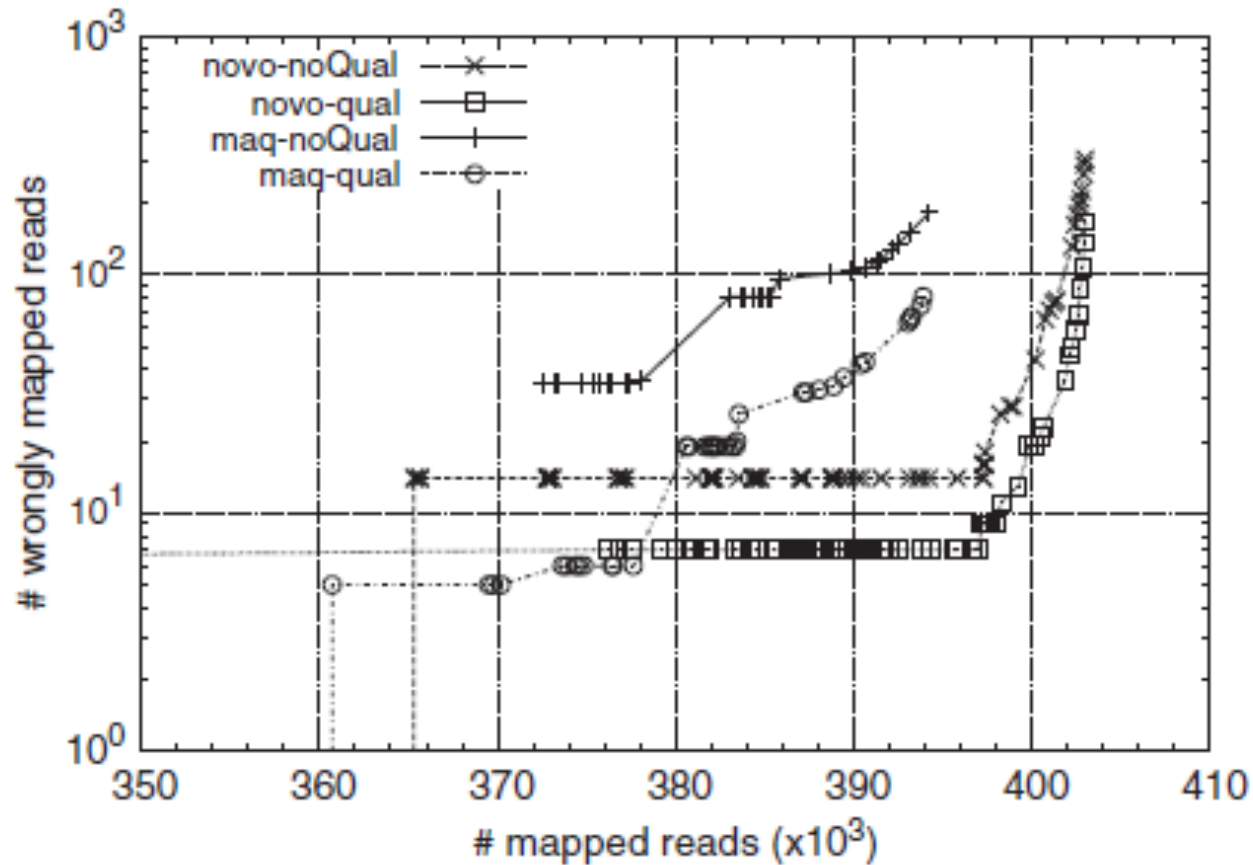
PCR duplicates

- **2nd generation sequencers are not single-molecule sequencers**
 - All have at least one PCR amplification step
 - Can result in duplicate DNA fragments
 - This can bias SNP calls or introduce false SNPs
- **Generally duplicates only make up a small fraction of the results**
 - Good libraries have < 2-3% of duplicates
 - SAMtools and Picard can identify and remove these when aligned against a reference genome
 - Debatable whether do this for RNA-seq and ChIP-seq
 - Depends on the complexity of the sample

PCR duplicates

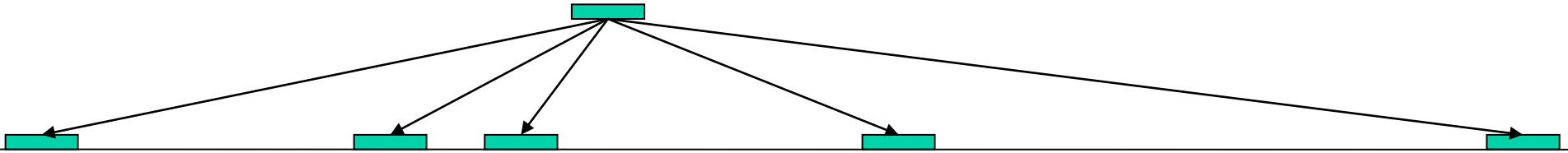
```
8661 8671 8681 8691 8701 8711 8721 8731 8741 8751 8761 8771 8781
901TCCCACTCTCAGACACTGAGAAAAGTGAGGCATGGGTTTCTGGGCTGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGTGAGGGAAAGGTGTAACCTGTTTGTGAGCCACAACATCT
M.....
AGCTCCCACTCTCAGACACTG tgggtttctgggctgggtacaggagctcgatgtgctttctctctacaggactgggtgagggaaagggtgtaacctgtttg
AGCTCCCACTCTCAGACACTG GTTTCTGGGCTGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGTGAGGGAAAGGTGTAACCTGTTTGTCA
AGCTCCCACTCTCAGACACTG GTTTCTGGGCTGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGTAAGGGAAAGGTGTAACCTGTTTGTCA
AGCTCCCACTCTCAGACACTG GTTTCTGGGCTGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGAGAGGGAAAGGTGTAACCTGTTTGTCA
AGCTCCCACTCTCAGACACTG GTTTCTGGGCTGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGTGAGGGAAAGGTGTAACCTGTTTGTCA
AGCTCCCACTCTCAGACACTG GTTTCTGGGCTGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGTAAGGGAAAGGTGTAACCTGTTTGTCA
AGCTCCCACTCTCAGACACTGAGAAAAGTGAGGCA GTTTCTGGGCTGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGAGAGGGAAAGGTGTAACCTGTTTGTCA
agctcccaactctcagacacttgagaaaagtgagggcatgggtttctggg CGATGTGCTTCTCTCTACAAGACTGGTGAGGGAAAGGTGTAACCTGTTTGTGAGCCACAACATCT
agctcccaactctctgacacttgagaaaagtgagggcatgggtttctggg tataacctatttgtagccacacacatct
agctcccaactctcagacacttgagaaaagtgagggcatgggtttctggg TAACCTGTTTGTGAGCCACAACATCT
agctcccaactctctgacacttgagaaaagtgagggcatgggtttctggg GTTTGTGAGCCACAACATCT
agctcccaactctcagacacttgagaaaagtgagggcatgggtttctggg GTTTGTGAGCCACAACATCT
agctcccaactctcagacacttgagaaaagtgagggcatgggtttctggg GTTTGTGAGCCACAACATCT
AACTGAGAAAAGTGAGGCATGGGTTTCTGGGCTGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGTGAGG GTTTGTGAGCCACAACATCT
AACTGAGAAAAGTGAGGCATGGGTTTATGGGATGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGTGAGG GTTTGTGAGCCACAACATCT
AACTGAGAAAAGTGAGGCATGGGTTTCTGGGCTGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGTGAGG
AACTGAGAAAAGTGAGGCATGGGTTTATGGGATGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGTGAGG
AACTGAGAAAAGTGAGGCATGGGTTTCTGGGCTGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGTGAGG
AACTGAGAAAAGTGAGGCATGGGTTTCTGGGCTGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGTGAGG
GTTTCTGGGCTGGTACAGGAGCTCGATGTGCTTCTCTCTACAAGACTGGTGAGTGAAAGGTTTAAATTTGTTTGTCT
```

Base quality impacts on read mapping



Heng Li & Nils Homer.
Sequence alignment
algorithms for next-
generation sequencing.
Briefings in
Bioinformatics. Vol 11.
No 5. 473 483, 2010

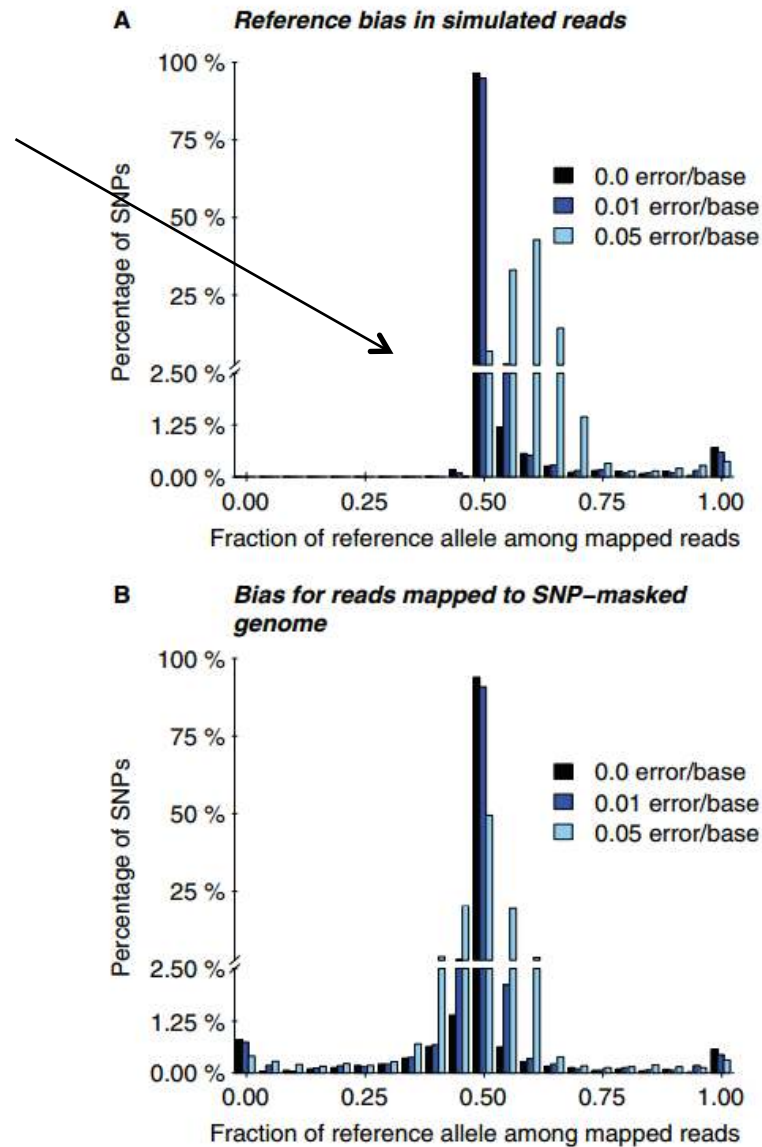
Multiple mapping reads



- A single read may occur more than once in the reference genome.
- Could be due to:
 - Paralogs (duplicated genes).
 - Transcripts which share exons.
 - Mutations in genotype relative to the reference.
 - Transposons and other common repetitive sequences
- Some aligners automatically assign a multi-mapping read to one of the locations at random (e.g. Tophat)
- Aligners may allow you to choose how these are dealt with – others may not

Allelic bias when SNP calling

Missing alternate allele



Methylation experiments

Unmethylated cytosine are converted to uracil

5' - atcgCCgataCga - 3'
3' - tagcgggCtatgct - 5'

Bisulfite
treatment

5' - atcgUUcgataUga - 3'

3' - tagcgggUtatgct - 5'

5' - atcgIIcgataIga - 3' (1)

3' - tagcAAgCtatAct - 5' (2)

Amplification

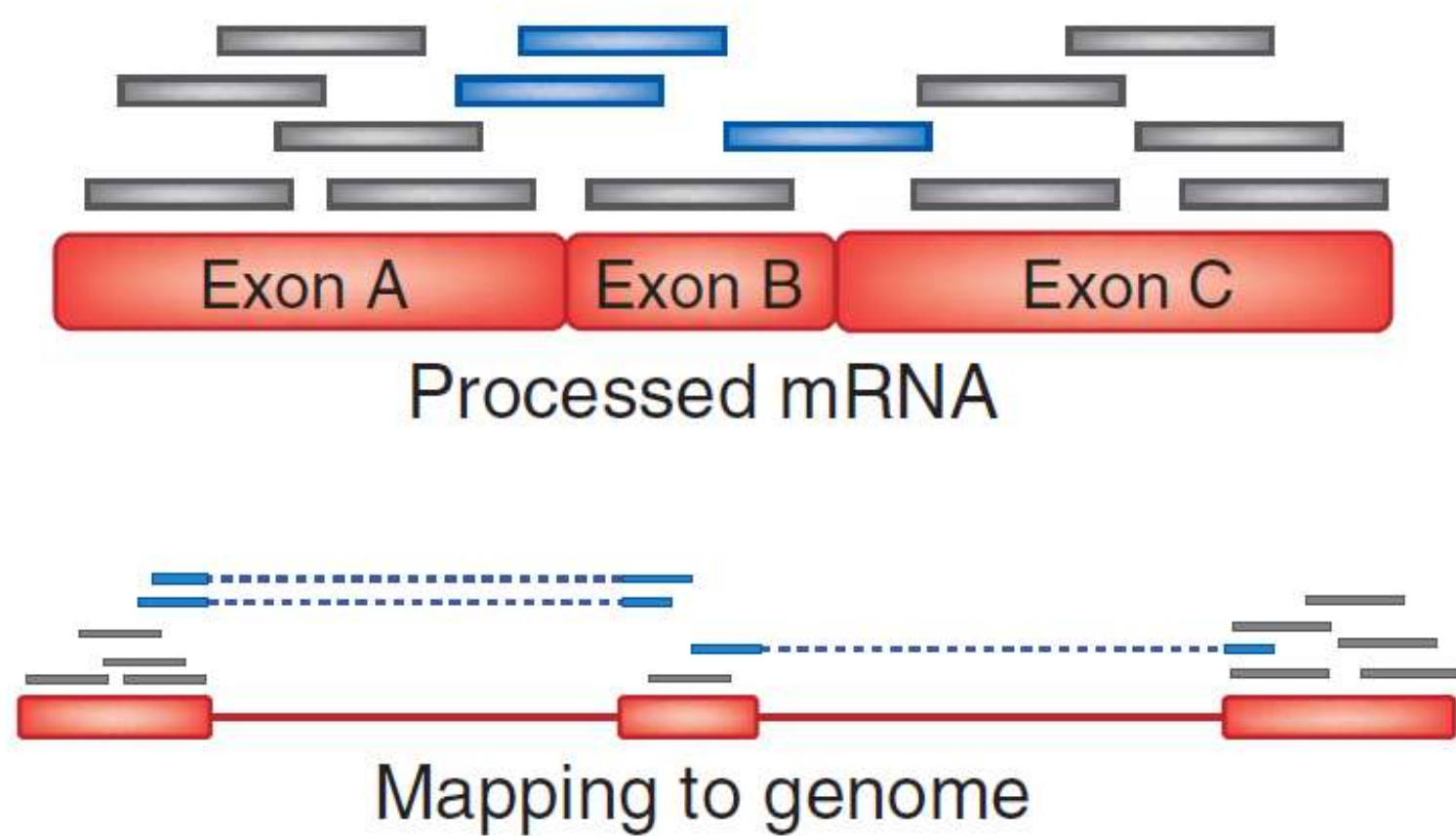
5' - atcgcccAatacga - 3' (3)

3' - tagcgggItatgct - 5' (4)

Methylation experiments

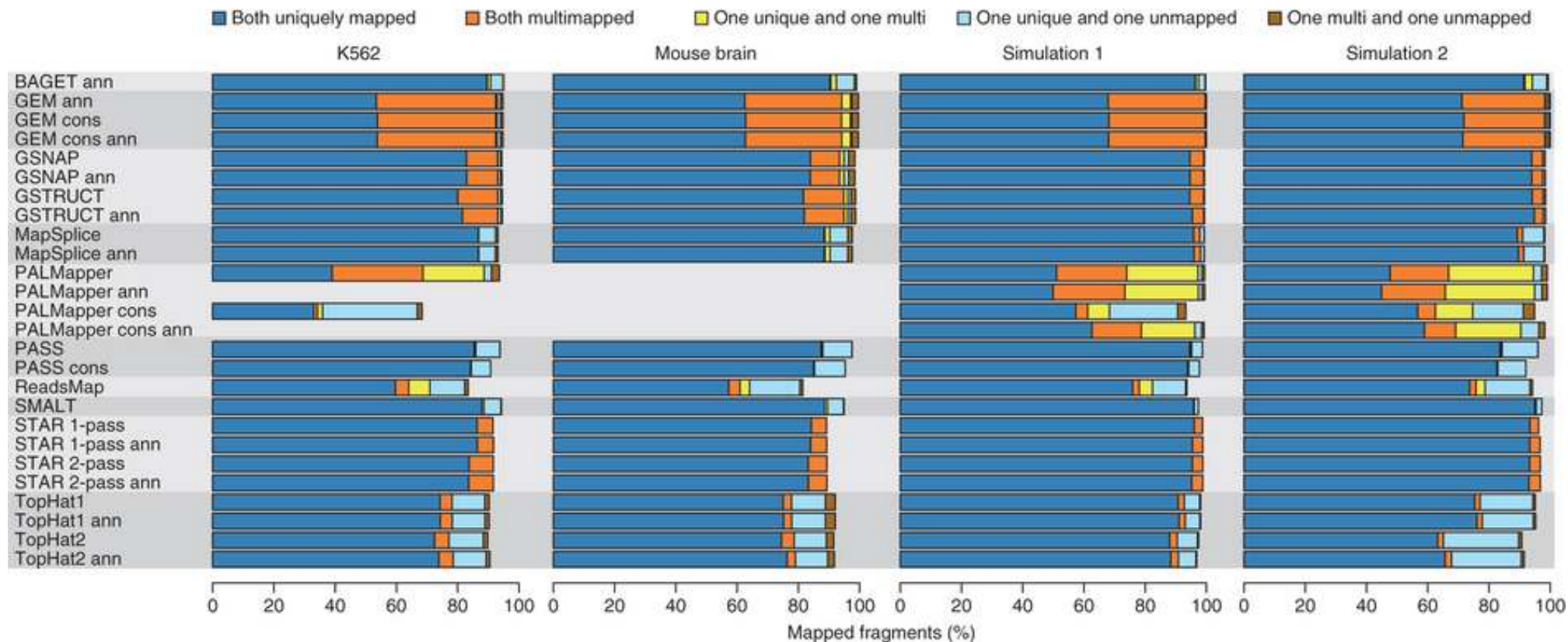
- Directly aligning reads against a reference will fail due to excessive mismatches in non-methylated regions
- Most packages deal with this by creating 2 reference sequences
 - One has all Cs converted to Ts
 - One has all Gs converted to As
- Convert Cs to Ts in all reads aligned against C->T reference
- Convert Gs to As in all reads aligned against G->A reference
- If there are no mutations or sequencing errors the reads will always map to one of the two references

Spliced-read mapping



- Need packages which can account for splice variants
- Examples: TopHat, STAR, GMAP, MapSplice

Spliced-read mapper evaluation



Local realignment to improve SNP/Indel detection

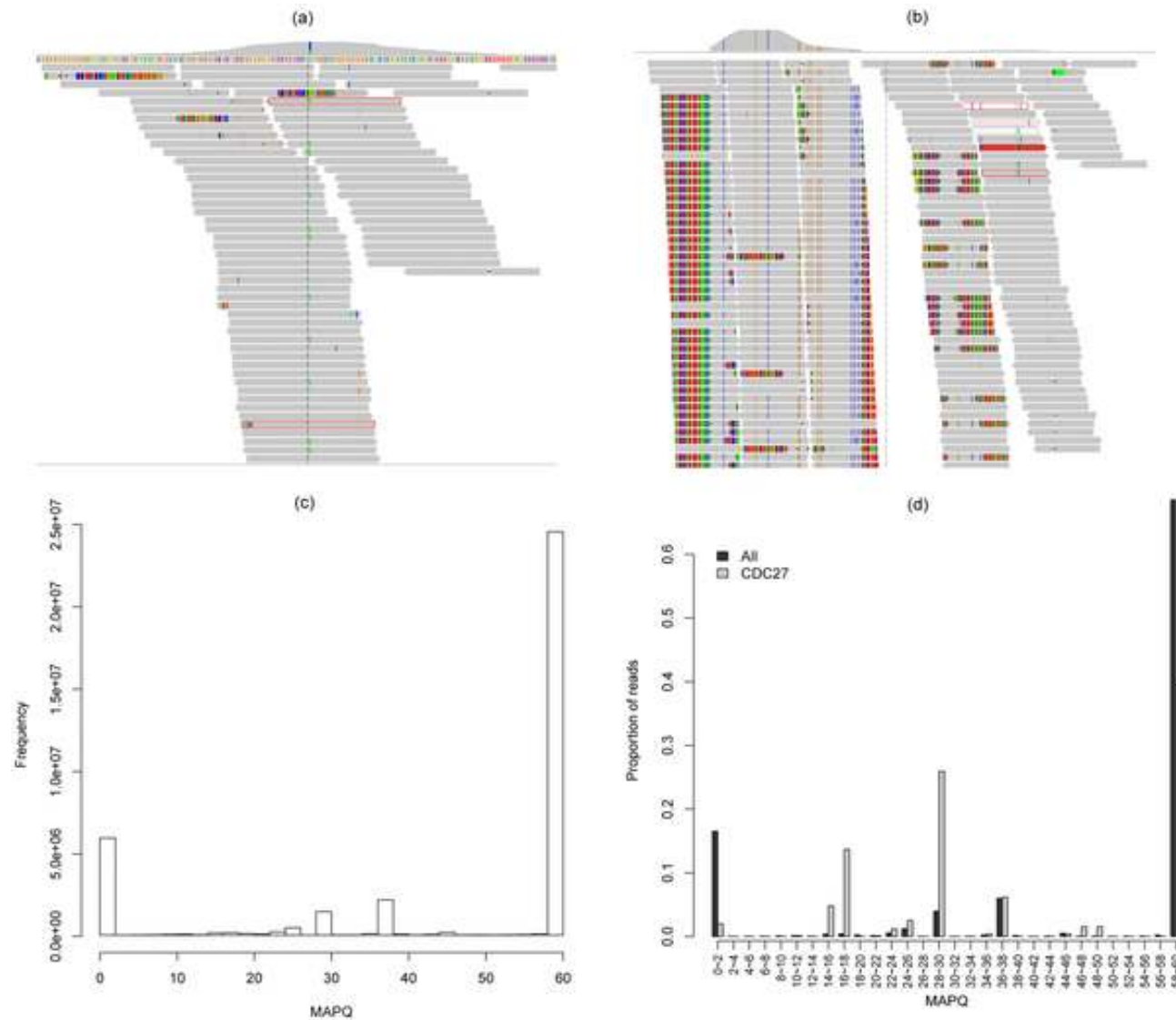
- Read aligners map each read (or read pair) independently of all other reads
- Around indels and other variants it can be helpful to make use of other metrics
 - e.g. Global median coverage for multi-mapping reads
- Tools such as GATK, SAMtools, Pindel and Breakdancer realign reads in the vicinity of variants to improve calls

http://www.broadinstitute.org/gsa/wiki/index.php/The_Genome_Analysis_Toolkit

Chen, K. BreakDancer: an algorithm for high-resolution mapping of genomic structural variation *Nature Methods* 6, 677 - 681 (2009)

Li H.*, Handsaker B.*, Wysoker A., Fennell T., Ruan J., Homer N., Marth G., Abecasis G., Durbin R. and 1000 Genome Project Data Processing Subgroup (2009) The Sequence alignment/map (SAM) format and SAMtools. *Bioinformatics*, 25, 2078-9

Figure 6. A visual examination of a spurious gene (CDC27).



All platforms have errors and artefacts



Illumina



PacBio



Roche 454



Ion Torrent

1. Removal of low quality bases
2. Removal of adaptor sequences
3. Platform specific artefacts (e.g homopolymers)

Table 2. Spurious genes having mutations detected in 30 samples.

CCDS ID	Gene symbol	Exon	# samples
CCDS11509.1	<i>CDC27</i>	13 th	36
CCDS12749.1	<i>CGB</i>	3 rd	36
CCDS12752.1	<i>CGB5</i>	1 st	36
CCDS41378.1	<i>NBPF11</i>	19 th	36
CCDS43407.1	<i>FAM153C</i>	4 th	36
CCDS5931.1	<i>MLL3</i>	42 nd	36
CCDS34703.1	<i>STAG3</i>	33 rd	34
CCDS5590.1	<i>POMZP3</i>	1 st	34
CCDS10638.1	<i>EIF3C</i>	8 th	32
CCDS30836.1	<i>NBPF14</i>	22 nd	31

CCDS: Consensus coding sequence. Exon: the specific exon in which the variants are detected.

doi:10.1371/journal.pone.0038470.t002

Illumina artefacts

Sequence-specific error profile of Illumina sequencers

Kensuke Nakamura^{1,*}, Taku Oshima², Takuya Morimoto^{2,3}, Shun Ikeda¹, Hirofumi Yoshikawa^{4,5}, Yuh Shiwa⁵, Shu Ishikawa², Margaret C. Linak⁶, Aki Hirai¹, Hiroki Takahashi¹, Md. Altaf-Ul-Amin¹, Naotake Ogasawara² and Shigehiko Kanaya¹

¹Graduate School of Information Science, ²Graduate School of Biological Sciences, Nara Institute of Science and Technology, 8916-5 Takayama-cho, Ikoma, Nara 630-0192, Japan, ³Biological Science Laboratories, Kao Corporation, 2606 Akabane, Ichikai, Haga, Tochigi 321-3497, ⁴Department of Bioscience, Tokyo University of Agriculture, ⁵Genome Research Center, NODAI Research Institute, Tokyo University of Agriculture, 1-1-1 Sakuragaoka Setagaya-ku, Tokyo, 156-8502, Japan and ⁶Department of Chemical Engineering and Material Science, University of Minnesota, 223 Amundson Hall, 421 Washington Avenue S.E., Minneapolis, MN 55455, USA

Received February 3, 2011; Revised April 25, 2011; Accepted April 26, 2011

ABSTRACT

We identified the sequence-specific starting positions of consecutive miscalls in the mapping of reads obtained from the Illumina Genome Analyser (GA). Detailed analysis of the miscall pattern indicated that the underlying mechanism involves sequence-specific interference of the base elongation process during sequencing. The two major sequence patterns that trigger this sequence-specific error (SSE) are: (i) inverted repeats and (ii) GGC sequences. We speculate that these sequences favor dephasing by inhibiting single-base

platforms [Illumina/Solexa Genome Analyser (4), Life Technologies/ABI SOLiD System (5) and Roche/454 Genome Sequencer FLX (6)], the Illumina Genome Analyser (GA) is, at the moment, the most popular choice for the analysis of genomic information (7). The Illumina/Solexa sequencers are characterized by: (i) solid-phase amplification and (ii) a cyclic reversible termination (CRT) process, also termed sequencing-by-synthesis (SBS) technology (8). The sequencer can generate hundreds of millions of relatively short (30–100 bp) read sequences per run.

The application of data obtained from this NGS technology can be roughly categorized into the following three

Nakamura, K. et al. Sequence-specific error profile of Illumina sequencers
Nucl. Acids Res. (2011) May 16, 2011

Illumina artefacts

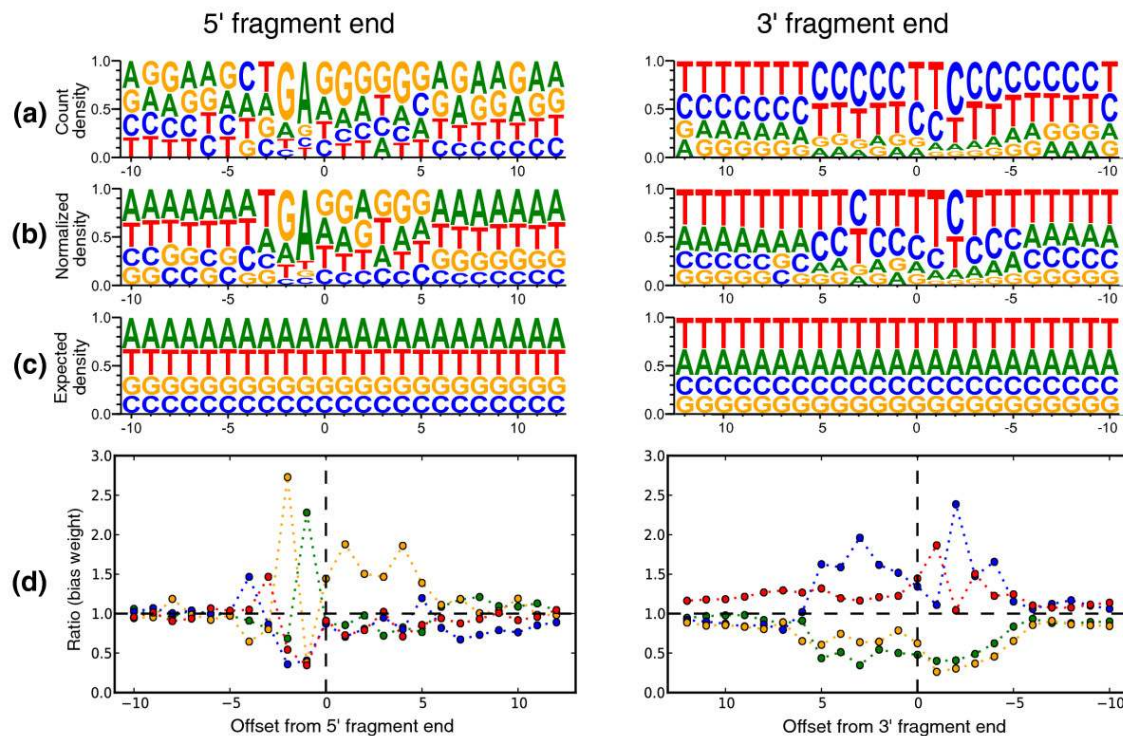
1. GC rich regions are under represented
 - a. PCR
 - b. Sequencing
2. Substitutions more common than insertions
3. GGC/GCC motif is associated with low quality and mismatches
4. Filtering low quality reads exacerbates low coverage of GC regions

Alignment software should ideally account for technology specific bias but generally does not

Your alignments are only as good as your library prep

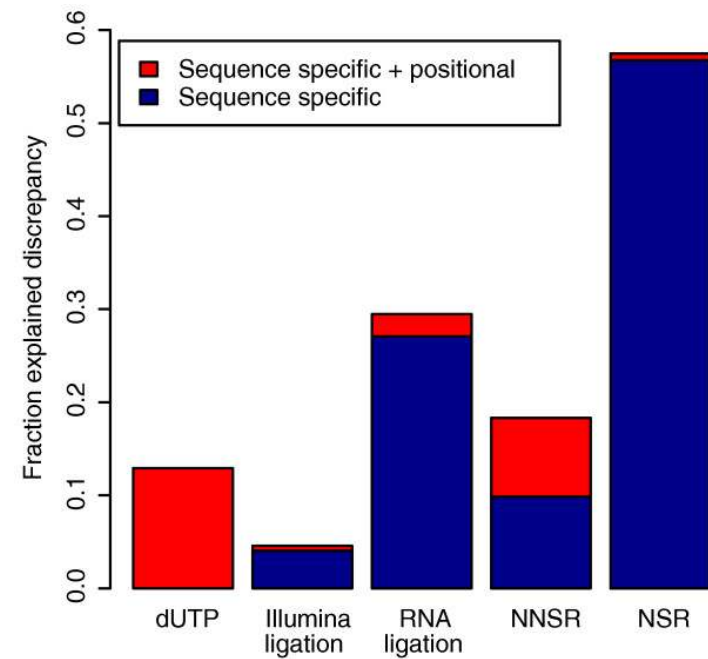
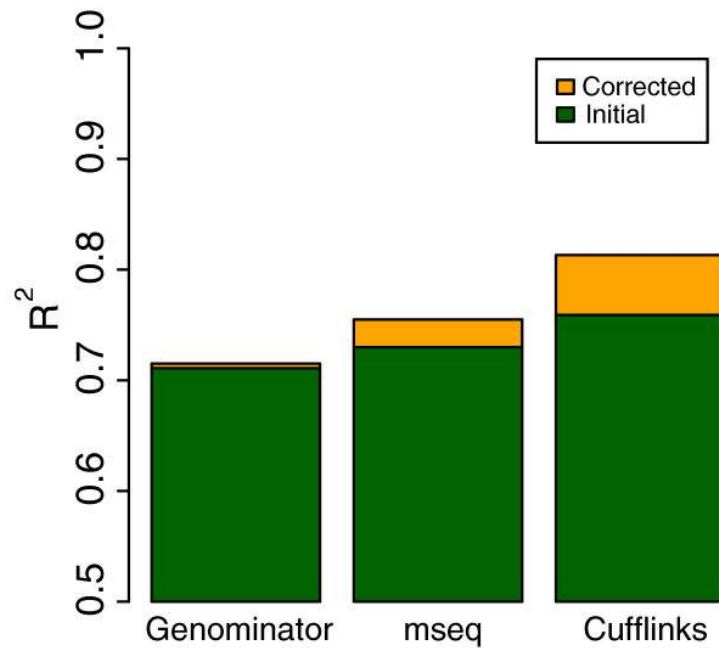
- Even if all other artefacts are removed:
- If your library prep is biased, your alignments will also reflect this bias

Tophat/Cufflinks aside



- Applies to random primed RNA-seq libraries
- Main potential biases:
 - Random hexamer priming biases
 - 5' or 3' ends of cDNA are likely to be mis-represented
 - Some packages correct for this (e.g. Tophat/Cufflinks)

Effect of bias correction

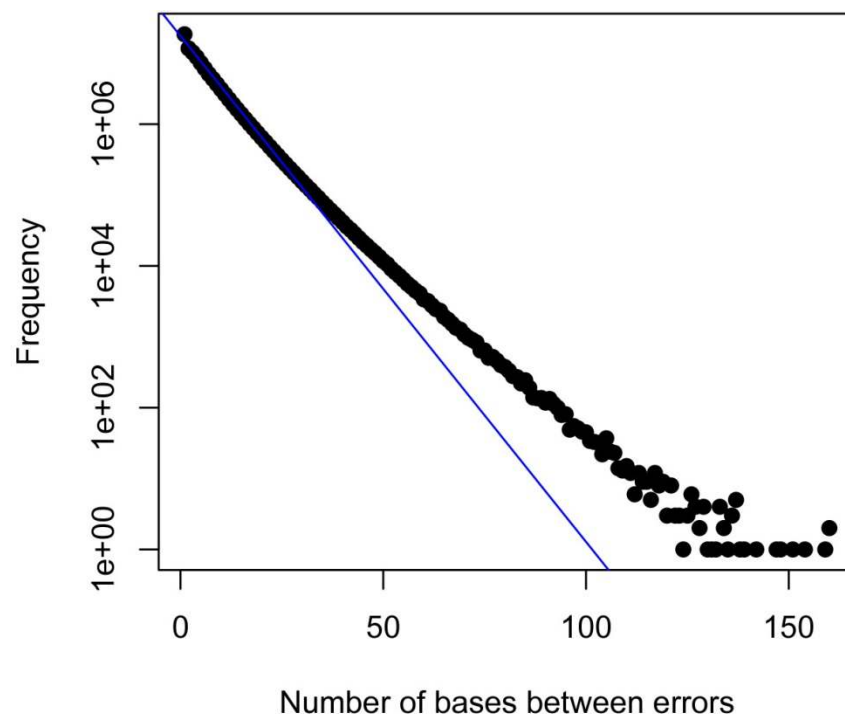


<http://genomebiology.com/2011/12/3/R22>

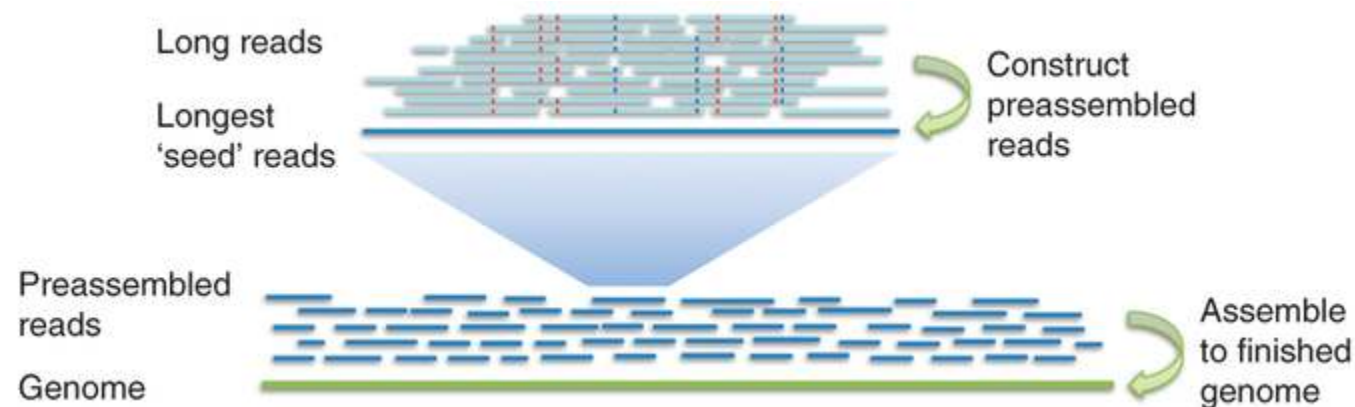
N.B. Out-dated version of Cufflinks used here

Pacific Biosciences alignment

- Median PacBio reads are 12kb with an single-read error rate of around 12-13%
- Most common distance between errors in a PacBio read is around 30bp
- Long length compensates and enables seeds to be located in many places to begin alignment
- However, this can be computationally costly
 - Need to balance number of seeds required to get a good alignment vs computational time
- PacBio developed BLASR to do this

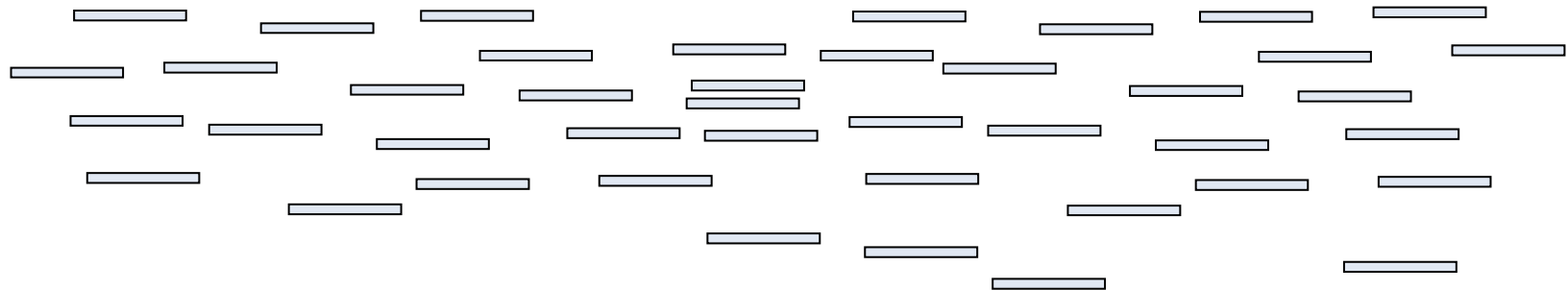
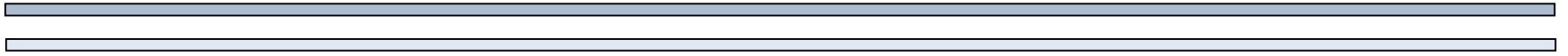


Pacific Biosciences read correction via alignment



- Alignment can also be used to generate a consensus to reduce the number of errors
- Most errors in PacBio seem to be randomly distributed
- If we have enough coverage, we can correct these errors by aligning all the short reads to the longer reads and correct based on the *consensus*
- These can then be used for denovo assembly
- Hierarchical Genome Assembly Process (HGAP)

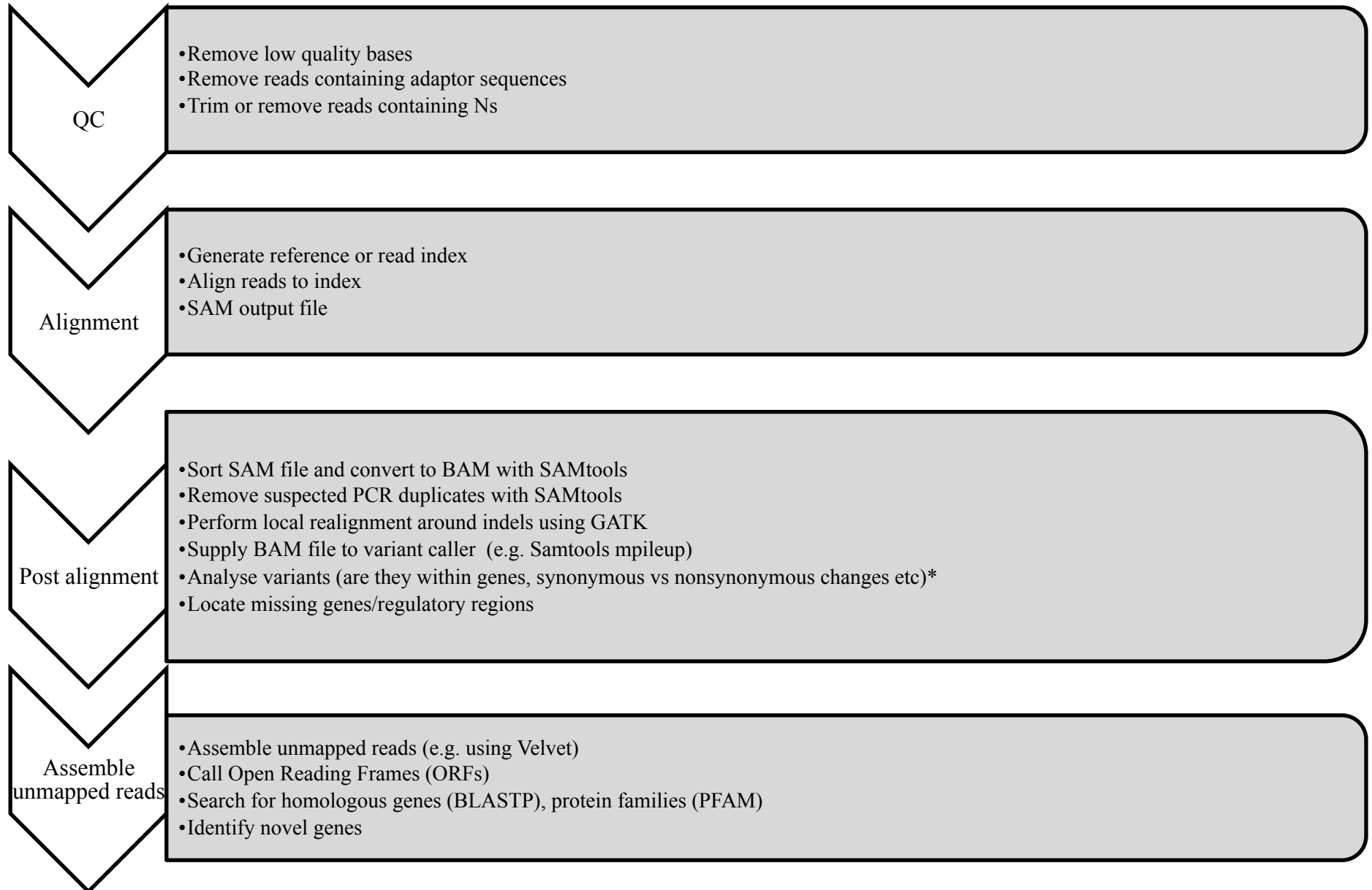
Unmapped reads



Unmapped reads

- Can be the result of:
 - Sequencing errors (should be small fraction if quality filtering applied before mapping)
 - Contamination
 - Excessive matches to repeats
 - Highly divergent regions between samples
 - Novel genetic material not present in reference
 - Plasmids
- Should be assembled de-novo with paired-end information if possible
- Resulting contigs run through MegaBlast against NCBI NT to check species
- Check against RepBase to remove repetitive contigs
- Call ORFs
- Blast ORFs using BlastP against NCBI NR or Swissprot and Blast2GO
- Run through PFAM

Typical alignment pipeline



* <http://bioinformatics.net.au/software.nesoni.shtml>

Contents

- **Alignment algorithms for short-reads**
 - Background – Blast (why can't we use it?)
 - Adapting hashed seed-extend algorithms to work with shorter reads
 - Indel detection
 - Suffix/Prefix Tries
 - Other alignment considerations
 - Typical alignment pipeline
 - **Haplotype methods of variant calling**

Haplotype SNP calling

- FreeBayes (<http://arxiv.org/pdf/1207.3907v2.pdf>)
- GATK – Haplotype caller
 - Haplotype calling in polyploids

ACCTGTA Reference Genome

Assume a SNP at both 5' A->T and 3' A->G in a diploid

Do we have a heterozygous?

Allele 1: ACCTGT**G**

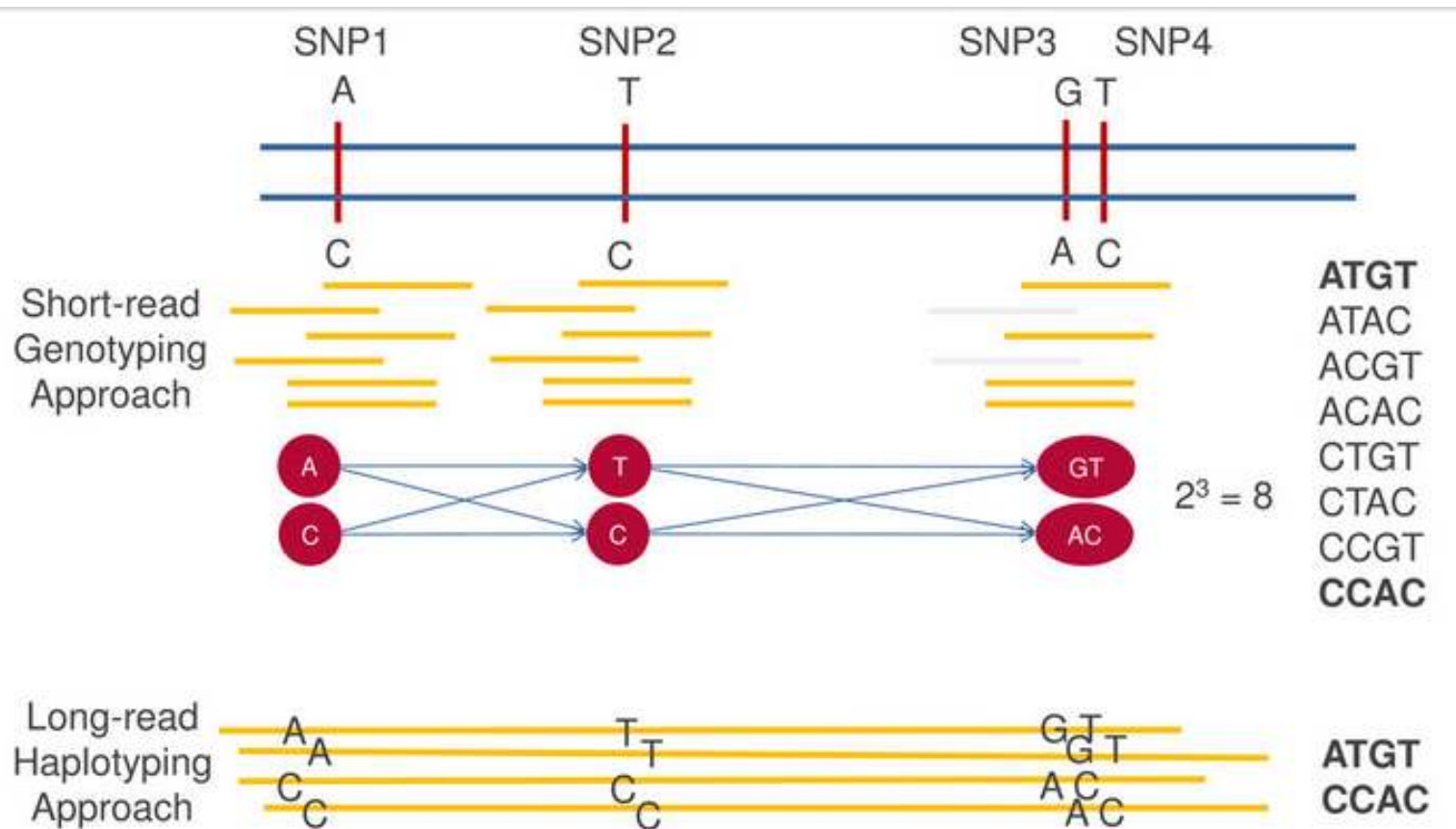
Allele 2: **T**CCTGT**C**

Or do we have a homozygous?

Allele 1: **T**CCTGT**G**

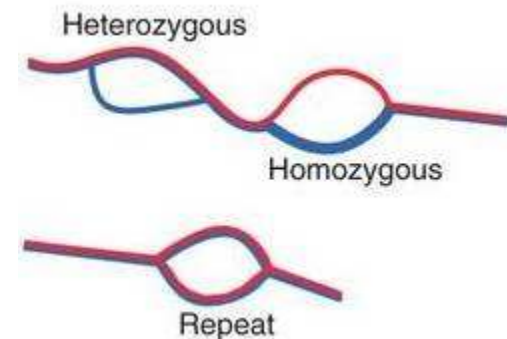
Allele 2: **T**CCTGT**G**

Haplotype issue calling – Long reads to the rescue



New methods of SNP calling

- Why align at all?
 - We only do this because of computational constraints
 - Ideally we want to assemble *denovo* and then align to reference genome
- Fermi and Cortex are tools to enable this:
 - *Denovo* genome assembler, but keeps track of differences which could be due to SNPs/Indels



Variant calling with de-novo assembly

Exploring single-sample SNP and INDEL calling with whole-genome de novo assembly

Heng Li^{1,*}

¹Broad Institute, 7 Cambridge Center, Cambridge, MA 02142, USA

Associate Editor: Dr. Michael Brudno

nature
genetics

ABSTRACT

Motivation: Eugene Myers in his stri suggested that in a string graph or e path spells a valid assembly. As a st every valid assembly of reads, such be constructed correctly, is in fact reads. In principle, every analysis bas sequencing (WGS) data, such as SNP calling, can also be achieved with uniti

De novo assembly and genotyping of variants using colored de Bruijn graphs

Zamin Iqbal^{1,2,5}, Mario Caccamo^{3,5}, Isaac Turner¹, Paul Flicek² & Gil McVean^{1,4}

Detecting genetic variants that are highly divergent from a reference sequence remains a major challenge in genome sequencing. We introduce *de novo* assembly algorithms using colored de Bruijn graphs for detecting and genotyping simple and complex genetic variants in an individual or population. We provide an efficient software implementation, Cortex, the first *de novo* assembler capable of assembling multiple eukaryotic genomes simultaneously. Four applications of Cortex are presented. First, we detect and validate both simple

a single suitable reference, as in ecological sequencing²¹. Fourth, methods for variant calling from mapped reads typically focus on a single variant type. However, in cases in which variants of different types cluster, focus on a single type can lead to errors, for example, through incorrect alignment around indel polymorphisms^{6,7}. Fifth, although there are methods for detecting large structural variants, such as using array comparative genomic hybridization (aCGH)²²⁻²⁵ and mapped reads^{11,12,14,26}, these cannot determine the exact location, size or allelic sequence of variants. Finally, mapping

Questions!