

Variant Calling

Resequencing-Based Genome Inference

Erik Garrison

University of Tennessee Health Science Center

Workshop on Genomics - Český Krumlov

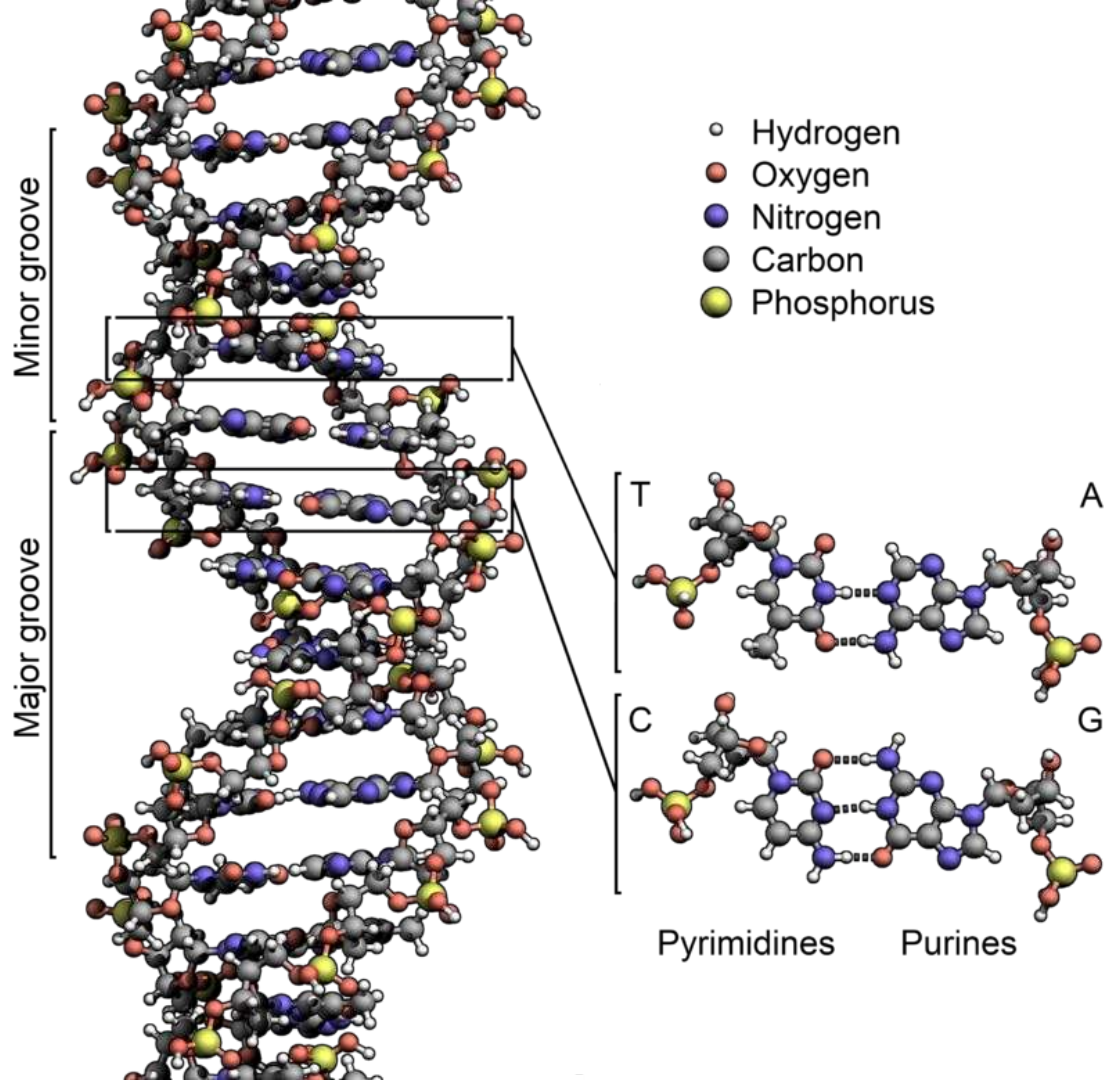
January 12, 2024

Overview

1. DNA variation and sequencing
2. Alignment to linear sequences
3. Error detection and genotyping
4. Learning to genotype
5. Practicals

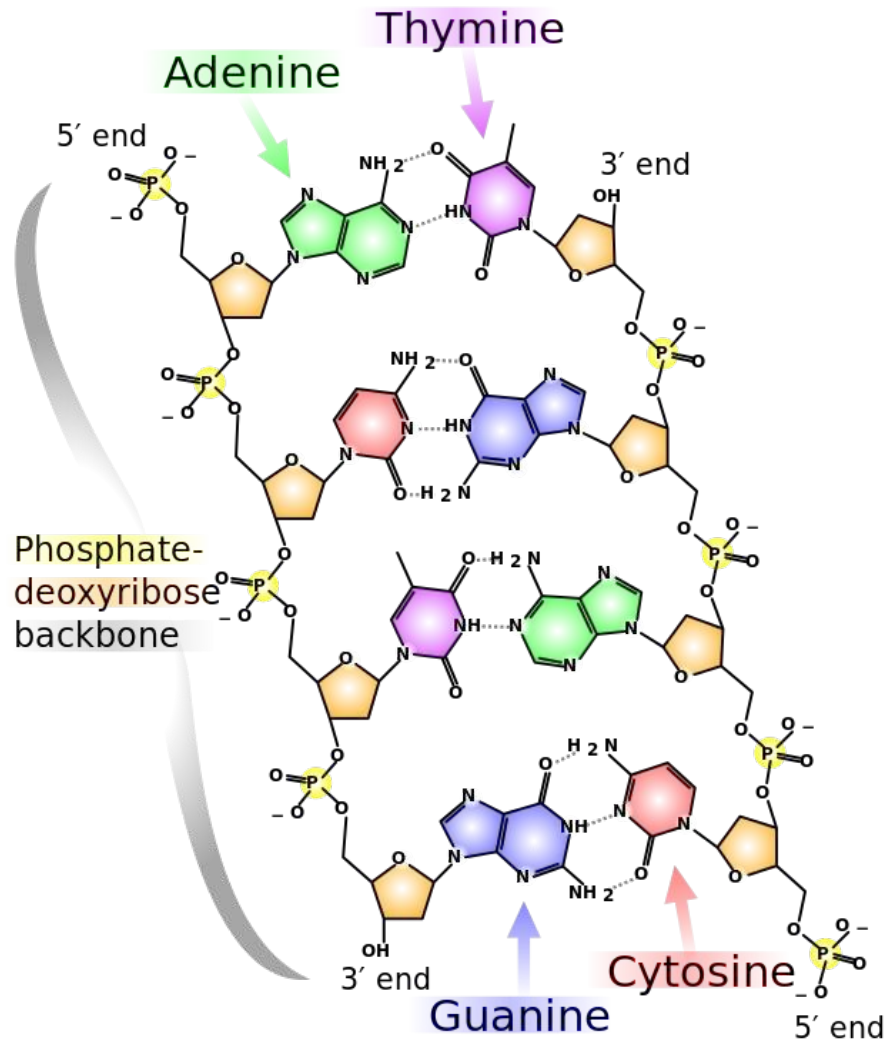
(small) Genomic variation

DNA



DNA

“twisted” view of ladder



A SNP

A point mutation in which one base is swapped for another.

AATTAGCCATTA

AATTAGTCATTA

An INDEL

A mutation that results from the gain or loss of sequence.

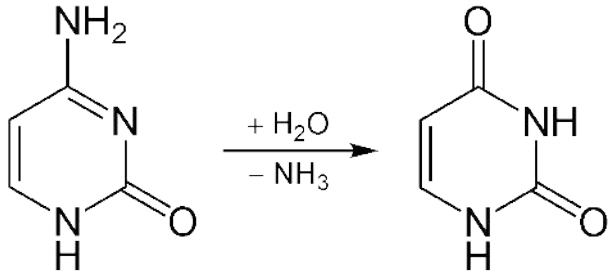
AATTAGCCATTA

AATTA--CATTA

(some) causes of SNPs

- Deamination

- cytosine → uracil
- 5-methylcytosine → thymine
- guanine → xanthine (mispairs to A-T bp)
- adenine → hypoxanthine (mispairs to G-C bp)



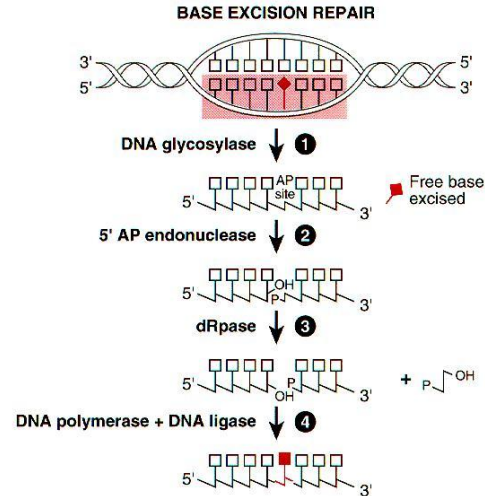
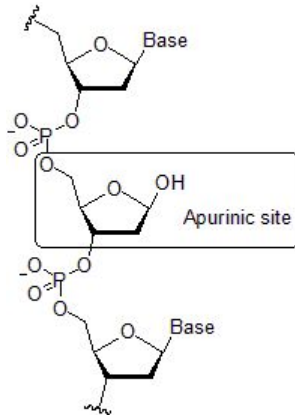
Deamination of Cytosine to Uracil

<http://en.wikipedia.org/wiki/Deamination>

(some) causes of SNPs

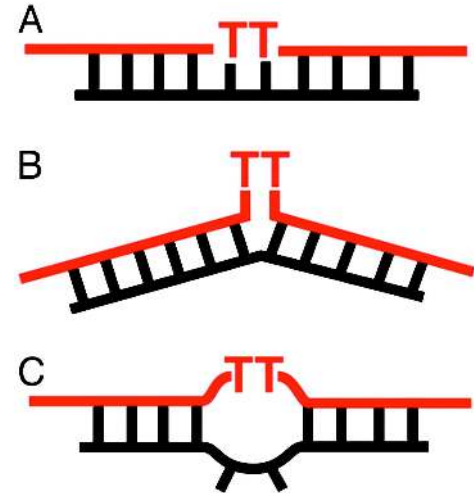
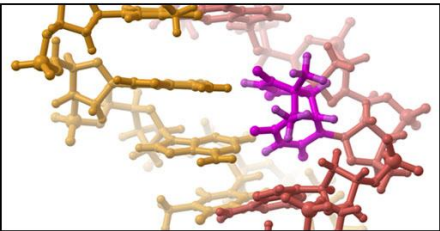
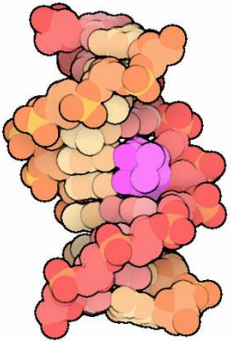
- Depurination

- purines are cleaved from DNA sugar backbone (5000/cell/day, pyrimidines at much lower rate)
- Base excision repair (BER) can fail → mutation



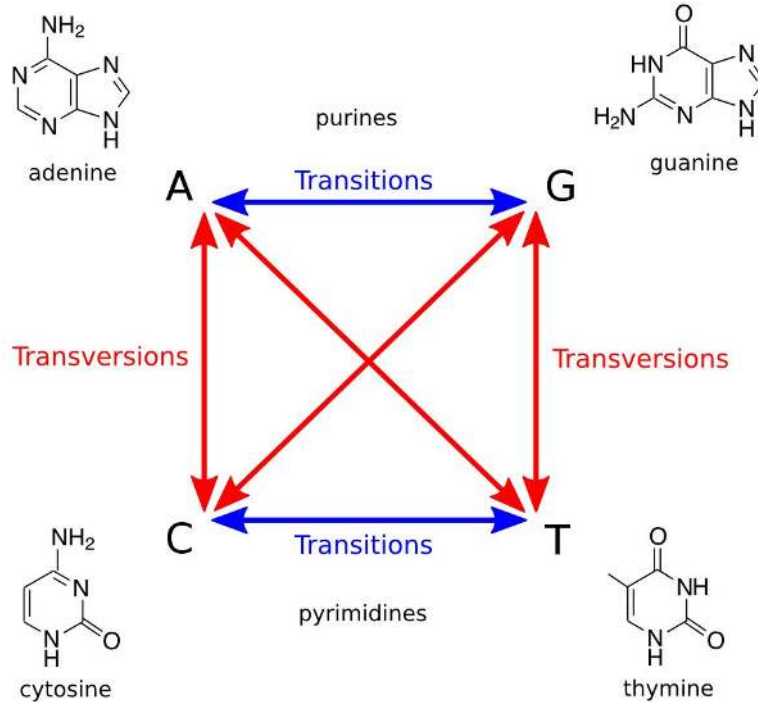
Multi-base events (MNPs)

- MNPs
 - thymine dimerization (UV induced)
 - other (e.g. oxidative stress induced)

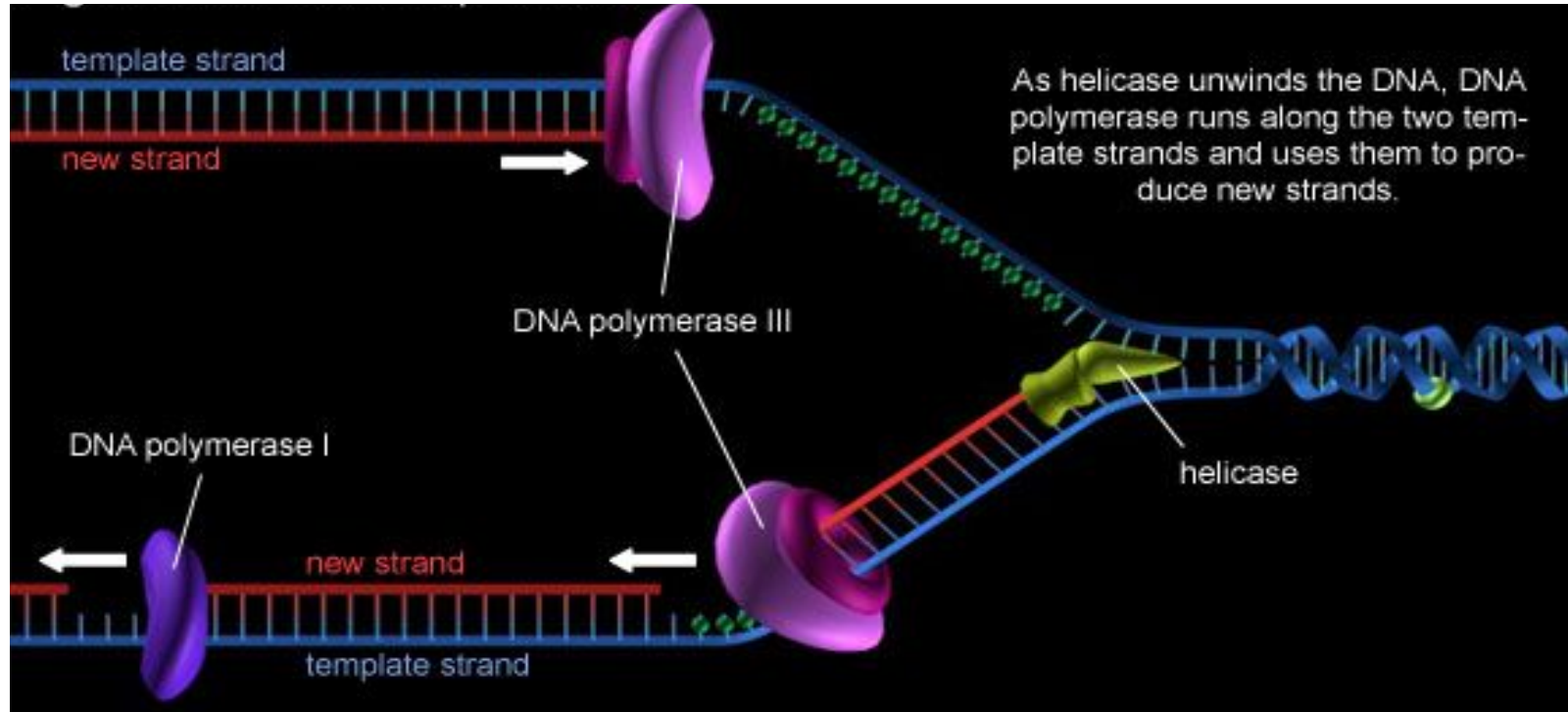


Transitions and transversions

In general **transitions** are **2-3 times more common than transversions**. (But this depends on biological context.)

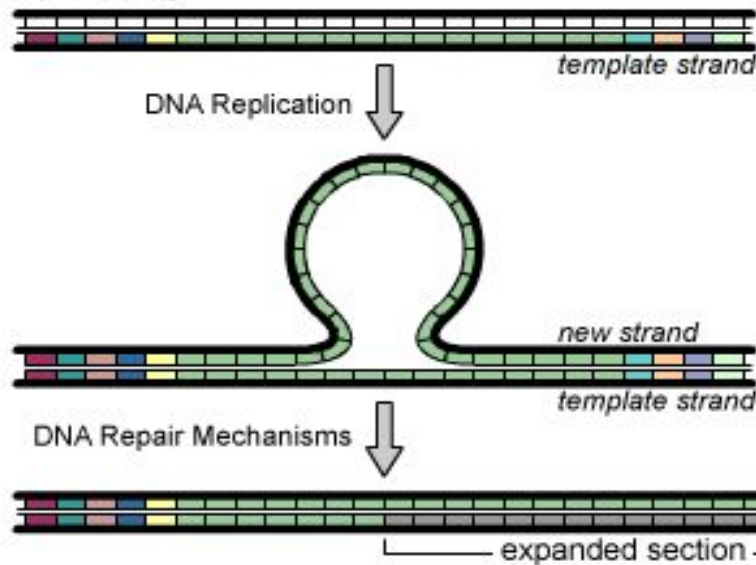


DNA replication



Polymerase *slippage*

A) Slippage Event



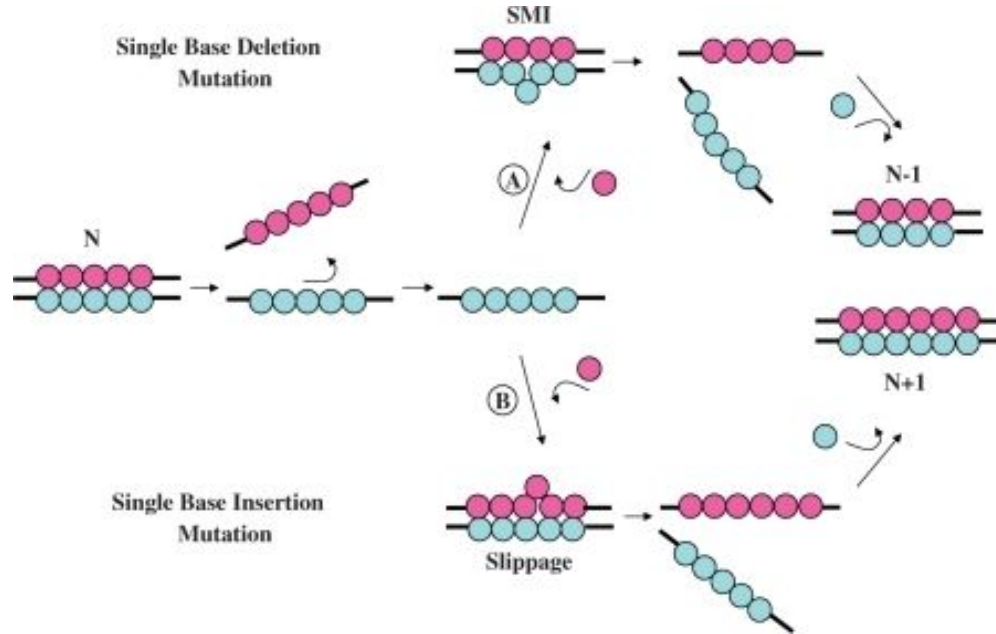
(A) During replication, polymerase slippage and subsequent reattachment may cause a bubble to form in the new strand. Slippage is thought to occur in sections of DNA with repeated patterns of bases (such as CAG), represented here by matching colors. Then, DNA repair mechanisms realign the template strand with the new strand and the bubble is straightened out. The resulting double helix is thus expanded.

B) No Slippage



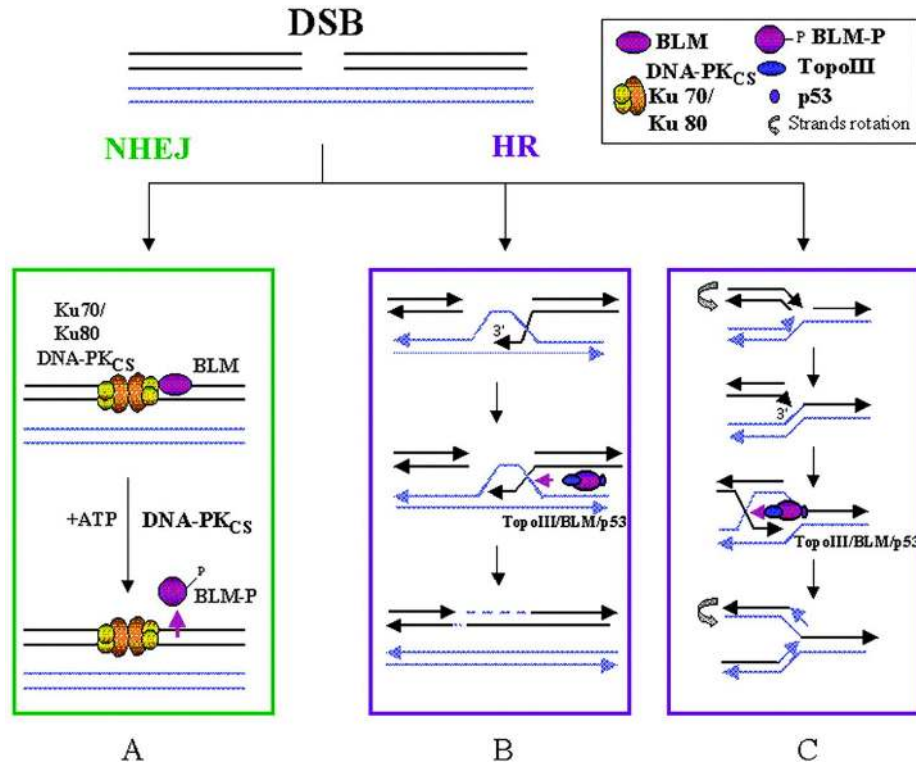
(B) Polymerase slippage, as theorized, cannot occur in DNA without repeating patterns of bases.

Insertions and deletions via slippage

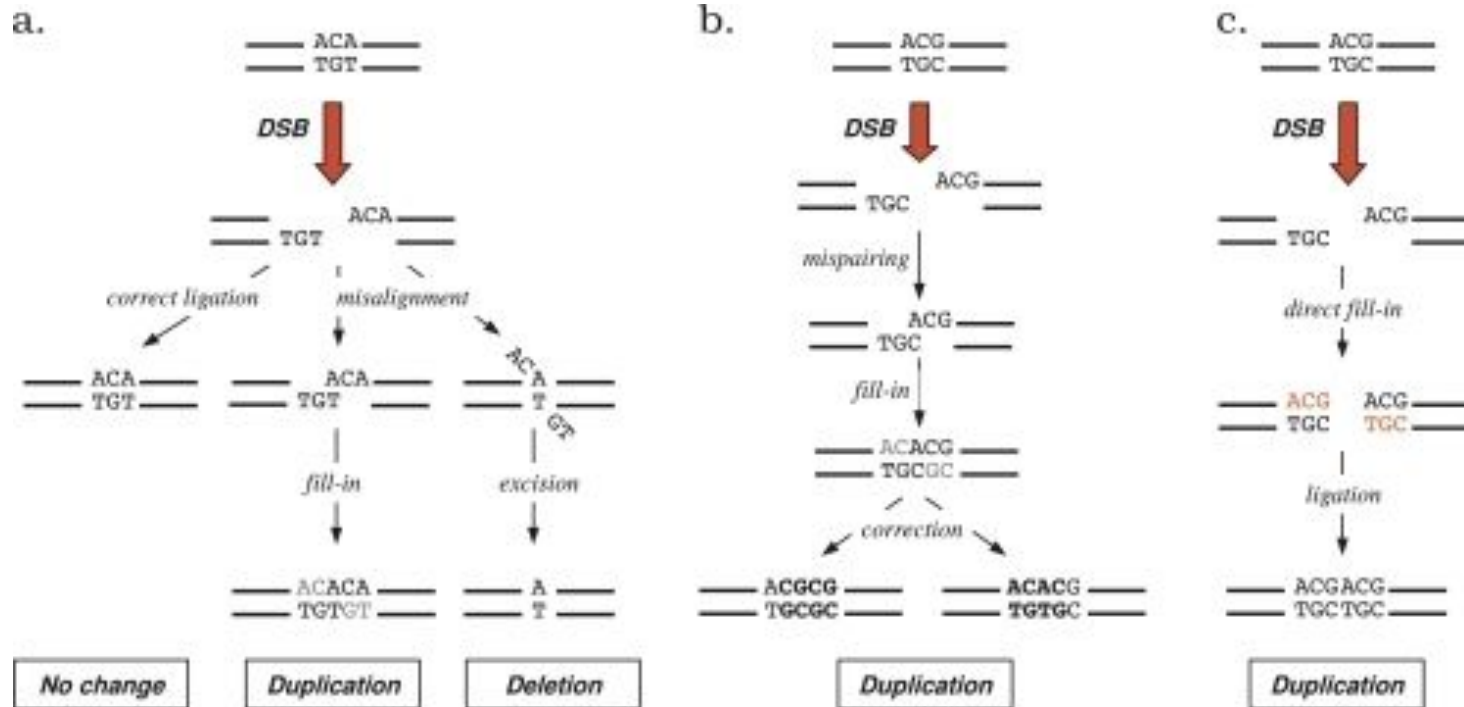


Energetic signatures of single base bulges: thermodynamic consequences and biological implications. Minetti CA, Remeta DP, Dickstein R, Breslauer K.J - Nucleic Acids Res. (2009)

Double-stranded break repair

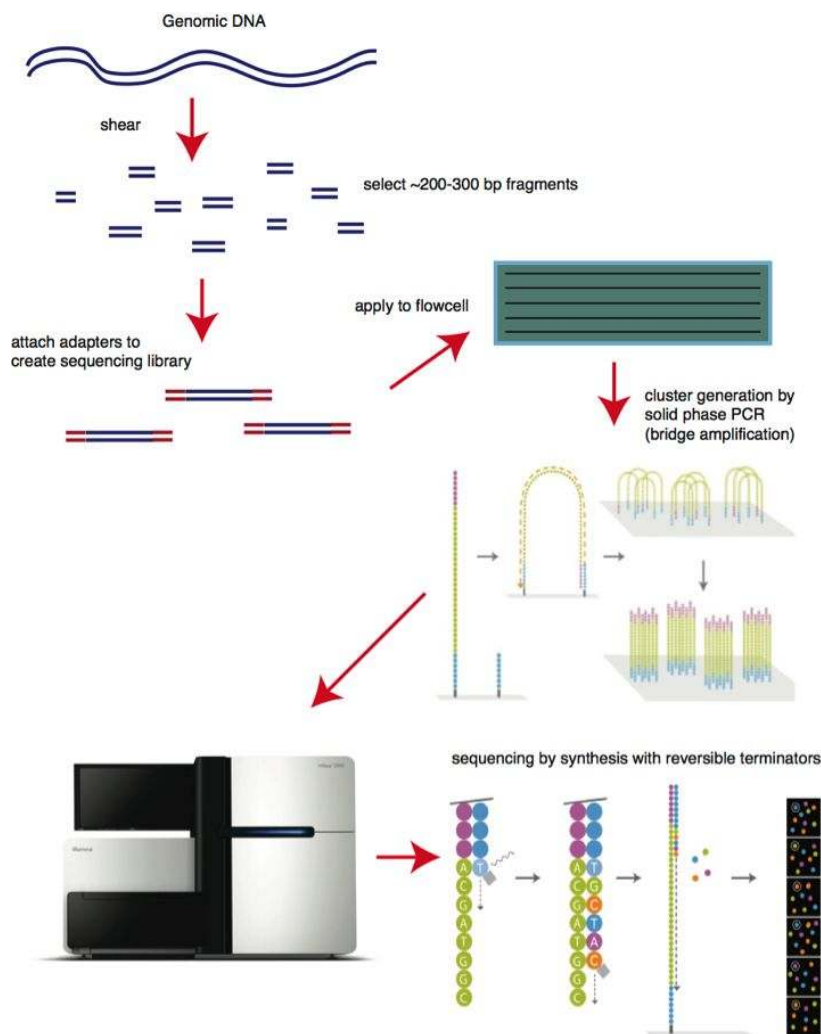


NHEJ-derived indels



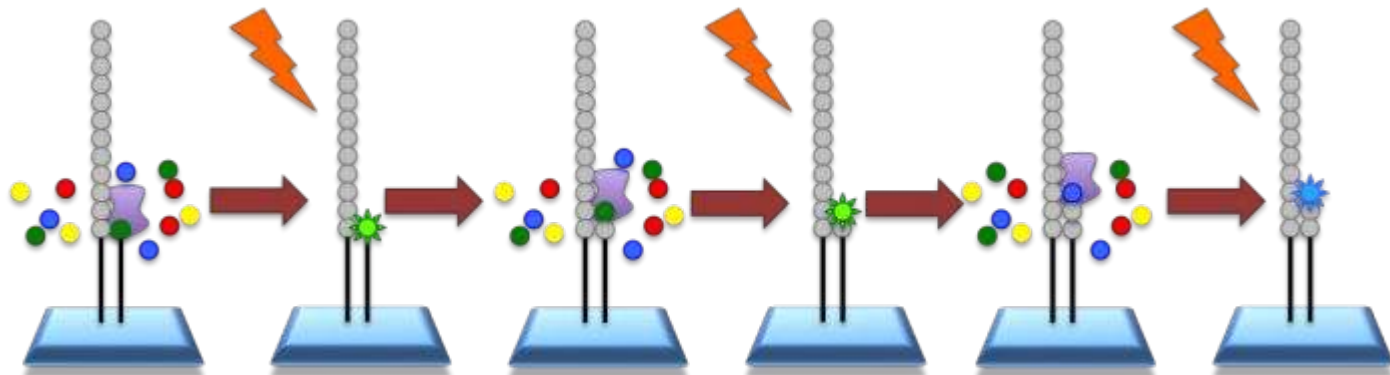
DNA Slippage Occurs at Microsatellite Loci without Minimal Threshold Length in Humans: A Comparative Genomic Approach. Leclercq S, Rivals E, Jarne P - Genome Biol Evol (2010)

Genome sequencing recap

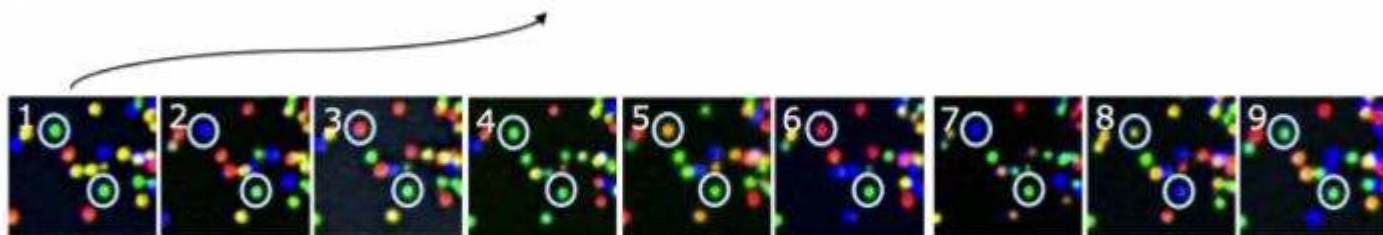


Sequencing by synthesis (Illumina)

n.b. Not exactly how this works on newer systems...



TGCTACGAT...



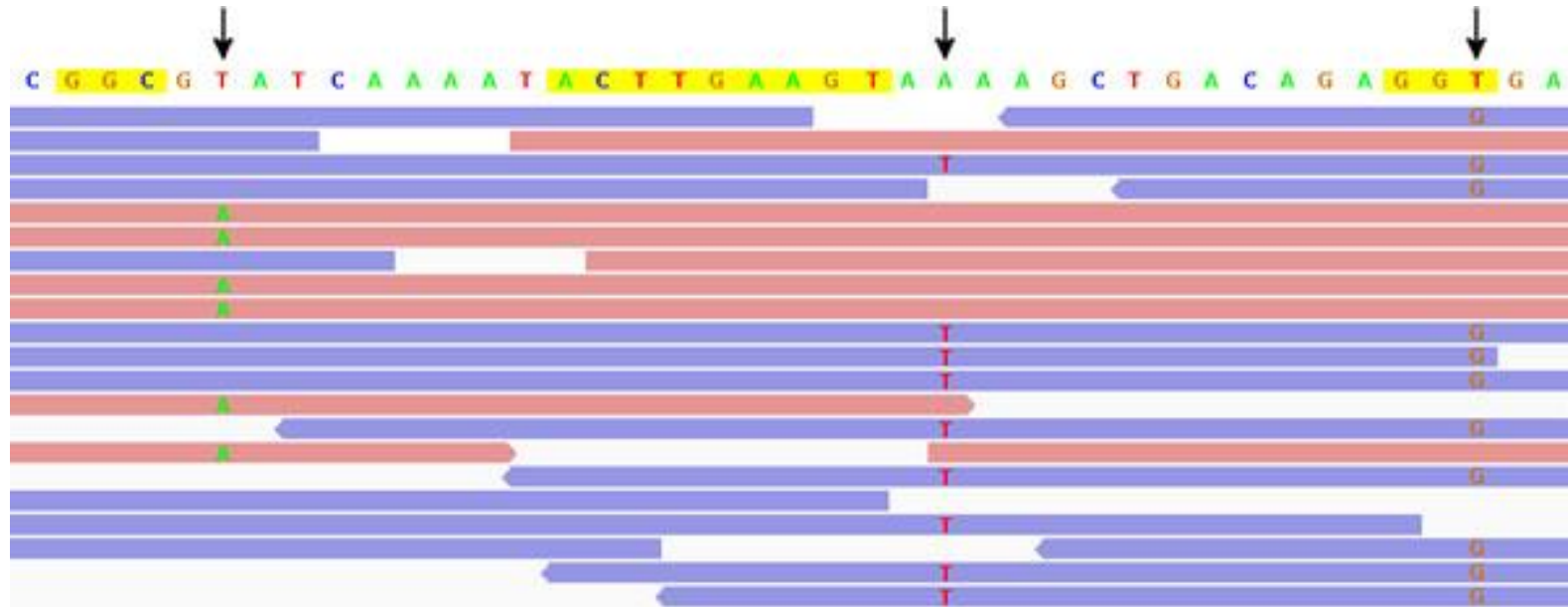
TTTTTTTGT...

What can go wrong?

1. Input artifacts, problems with library prep
 - a. replication in PCR has no error-correction (→ SNPs)
 - b. no quaternary structures (e.g. clamp) to prevent slippage (→ indels)
 - c. chimeras...
 - d. duplicates (worse if they are errors)
2. Sequencing-by-synthesis
 - a. phasing of step
 - i. synthesis reaction efficiency is not 100%
 - ii. particularly bad in A/T homopolymers
 - b. certain *context specific errors*
 - i. vary by sequencing protocol, device
 - ii. often strand-specific

Example: Context specific errors

Show up as strand-specific errors:



Context specific errors (motifs)

Rank	Context	FER [%]	RER [%]	ERD [%]
1	ACGGCGGT	26.1	0.5	25.6
2	GTGGCGGT	25.1	0.7	24.4
3	GCGGCGGT	22.9	0.7	22.2
4	GTGGCTGT	22.4	0.6	21.8
5	ATGGCGGT	21.2	1.0	20.3
6	NCGGCGGT	20.0	0.7	19.3
7	GTGGCTTG	20.2	1.2	19.0
8	GNGGCGGT	19.2	0.7	18.5
9	GCGGCTGT	18.8	0.7	18.1
10	ACGGCTGT	18.6	0.8	17.7

← forward and reverse error rates for the ten most-common CSEs on a variety of illumina systems (in 2013)

Often GC-rich!

Changes in chemistry mean that this may not be such a big deal now, but this example is something to keep in mind!

Long read technologies

“3rd-gen” sequencing.

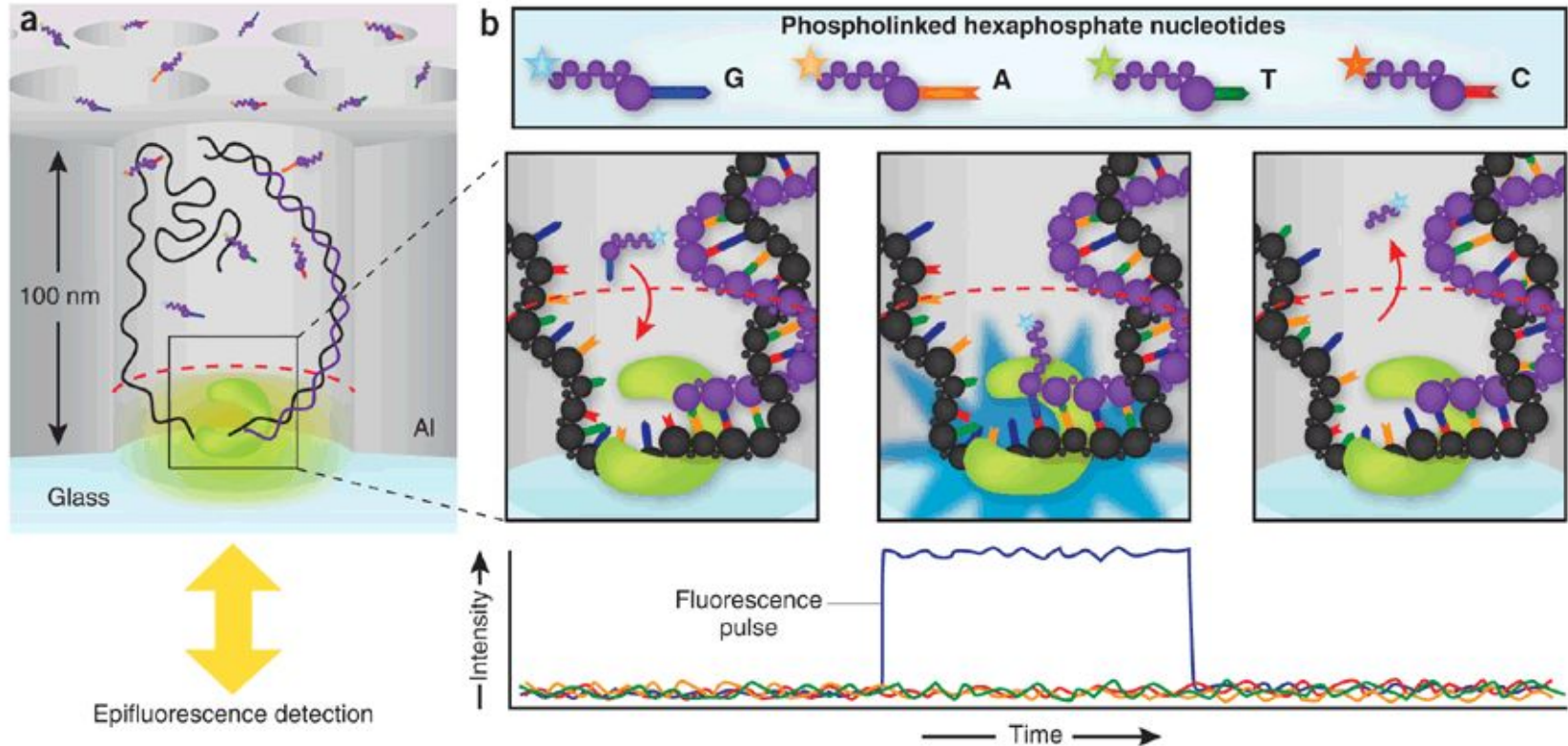
- Read single molecules (long too!)
- Have high error rates (10-15%)

*Pacific
Biosciences*



Oxford Nanopore

Pacific Biosciences sequencing

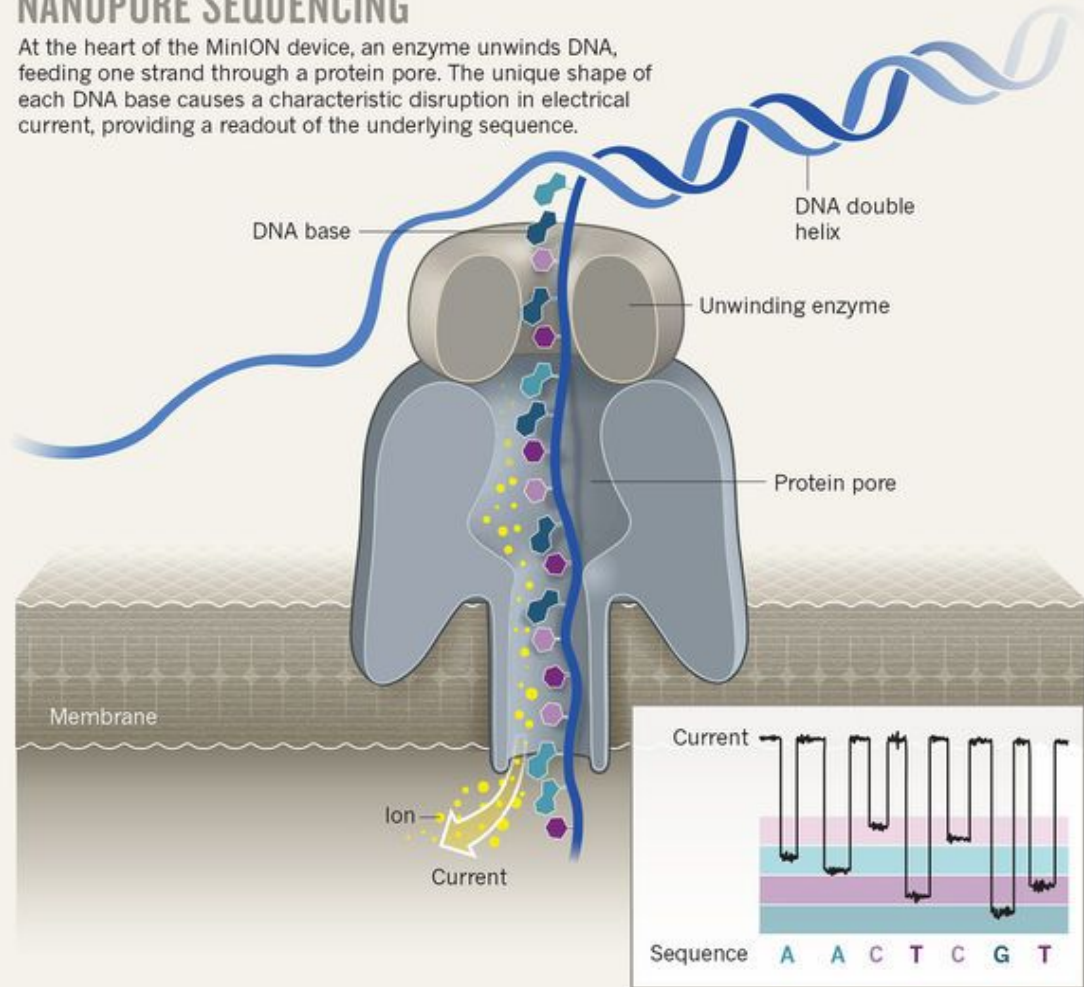


Oxford Nanopore sequencing



NANOPORE SEQUENCING

At the heart of the MinION device, an enzyme unwinds DNA, feeding one strand through a protein pore. The unique shape of each DNA base causes a characteristic disruption in electrical current, providing a readout of the underlying sequence.



Alignment is interpretation

Alignment

Covered in our last session!

Key idea for variant calling:

alignment provides a kind of interpretation of variation.

Changes in parameters, alignments, or sequence context can lead to changes in called alleles.

Seeing variation

These sequences have mutations between them.



CAAATAAGGAAATTTTCTGGAGTTCTATTATA
CAAATAAGGTTTGTATCTAGGTTATTATA

The image displays two DNA sequences, one above the other. Each nucleotide is represented by a colored letter: C (blue), A (green), A (red), T (yellow), A (green), G (red), G (yellow), A (green), A (red), T (yellow), T (red), T (blue), C (yellow), T (green), G (red), G (yellow), A (green), G (red), T (blue), T (green), C (red), T (yellow), A (green), T (red), T (blue), A (green), T (red), A (yellow), T (green), A (red), T (blue), A (green). The sequences are aligned to show mutations between them.

They are homologous but it's not easy to see.

Pairwise alignment

One solution, assuming a particular set of alignment parameters, has 3 indels and a SNP:

```
CAAATAAGGAAATTT - - - TCTGGAGTTCTATTATA  
CAAATAAGG - - - TTTGCTATCT - - AGGT - TATTATA
```

But if we use a higher gap-open penalty, things look different:

```
CAAATAAGGAAATTT - - TCTGGAGTTCTATTATA  
CAAATAAGG - - - TTTGCTATCTAGGT - TATTATA
```

Alignment = interpretation

Different parameterizations can yield different results.

Different results suggest “different” variation.

What kind of problems can this cause? (And how can we mitigate these issues?)

INDELs have multiple representations and require normalization for standard calling

Left alignment allows us to ensure that our representation is consistent across alignments and also variant calls.



CGTATGATCTAGCGCGCTAGCTAGCTAGC ← Left
CGTATGATCTA - - GCGCTAGCTAGCTAGC aligned

CGTATGATCTAGCGCGCTAGCTAGCTAGC
CGTATGATCTAGC - - GCTAGCTAGCTAGC

CGTATGATCTAGCGCGCTAGCTAGCTAGC
CGTATGATCTAGCGC - - TAGCTAGCTAGC

example: 1000G Phasel low coverage

chr15:81551110, ref:CTCTC alt:ATATA

ref: TGTCACTCGCTCTCTCTCTCTCTCTATATATATATTTGTGCAT
alt: TGTCACTCGCTCTCTCTCTCTATATATATATATATATTTGTGCAT

The diagram illustrates a sequence alignment between a reference (ref) and an alternative (alt) sequence. Both sequences are 28 nucleotides long. The reference sequence is TGTCACTCGCTCTCTCTCTCTCTCTATATATATATTTGTGCAT, and the alternative sequence is TGTCACTCGCTCTCTCTCTCTATATATATATATATATTTGTGCAT. The sequences are color-coded: A is red, C is blue, G is orange, and T is green. A blue bracket above the alignment spans from the first 'A' to the last 'T' of the common segment. Another blue bracket below the alignment spans from the first 'A' to the last 'T' of the common segment.

Interpreted as 3 SNPs

[illegible]

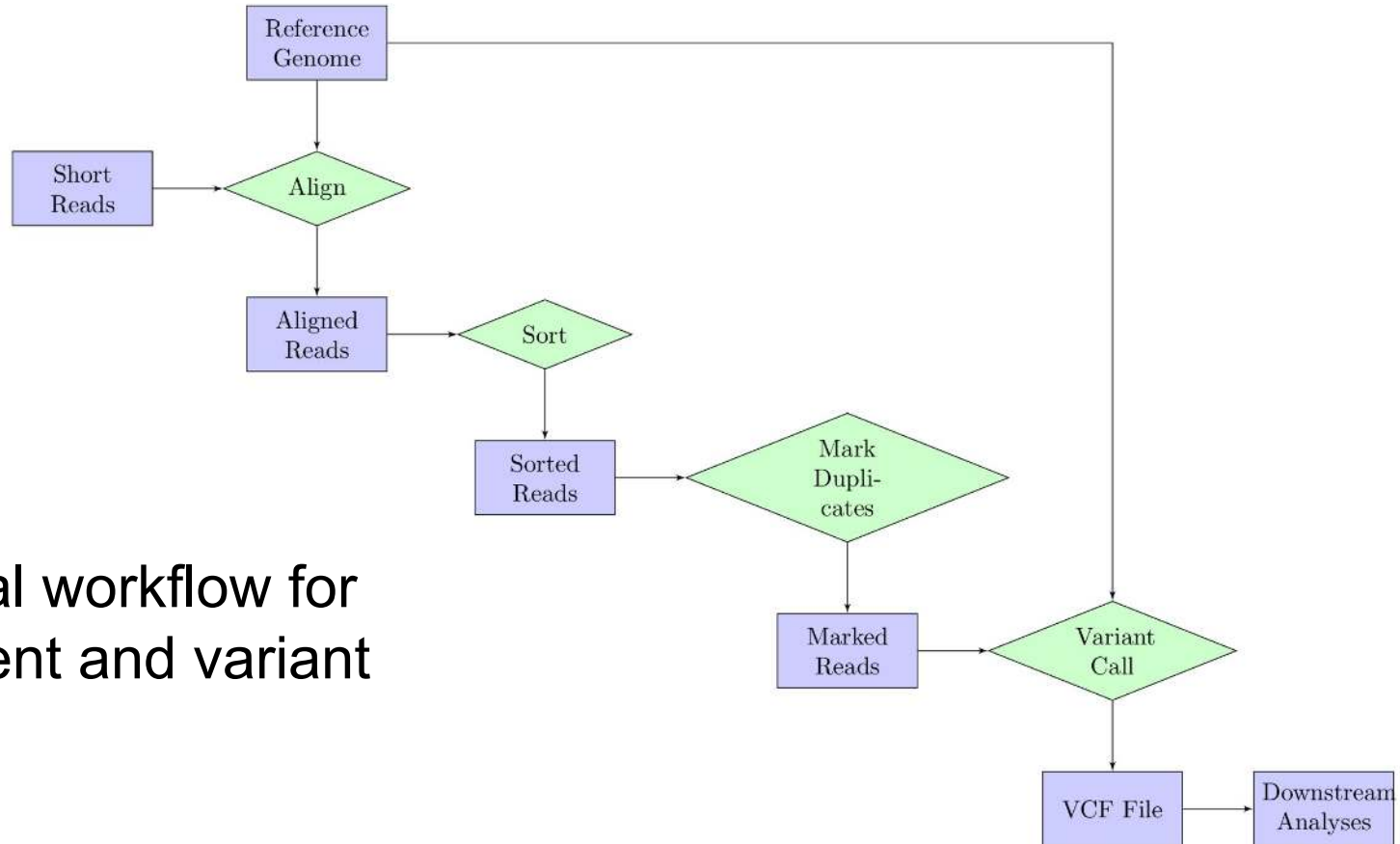
Interpreted as microsatellite expansion/contraction

example: 1000G Phasel low coverage
chr20:708257, ref:AGC alt:CGA

ref: TATAGAGAGAGAGAGAGAGAGC GAGAGAGAGAGAGAGAGAGGGAGAGACGGAGTT
alt: TATAGAGAGAGAGAGAGAGAGC GAGAGAGAGAGAGAGAGAGGGAGAGACGGAGTT

ref: TATAGAGAGAGAGAGAGAGAGC - - GAGAGAGAGAGAGAGAGAGGGAGAGACGGAGTT
alt: TATAGAGAGAGAGAGAGAG - - CGAGAGAGAGAGAGAGAGAGGGAGAGACGGAGTT

Processing alignments

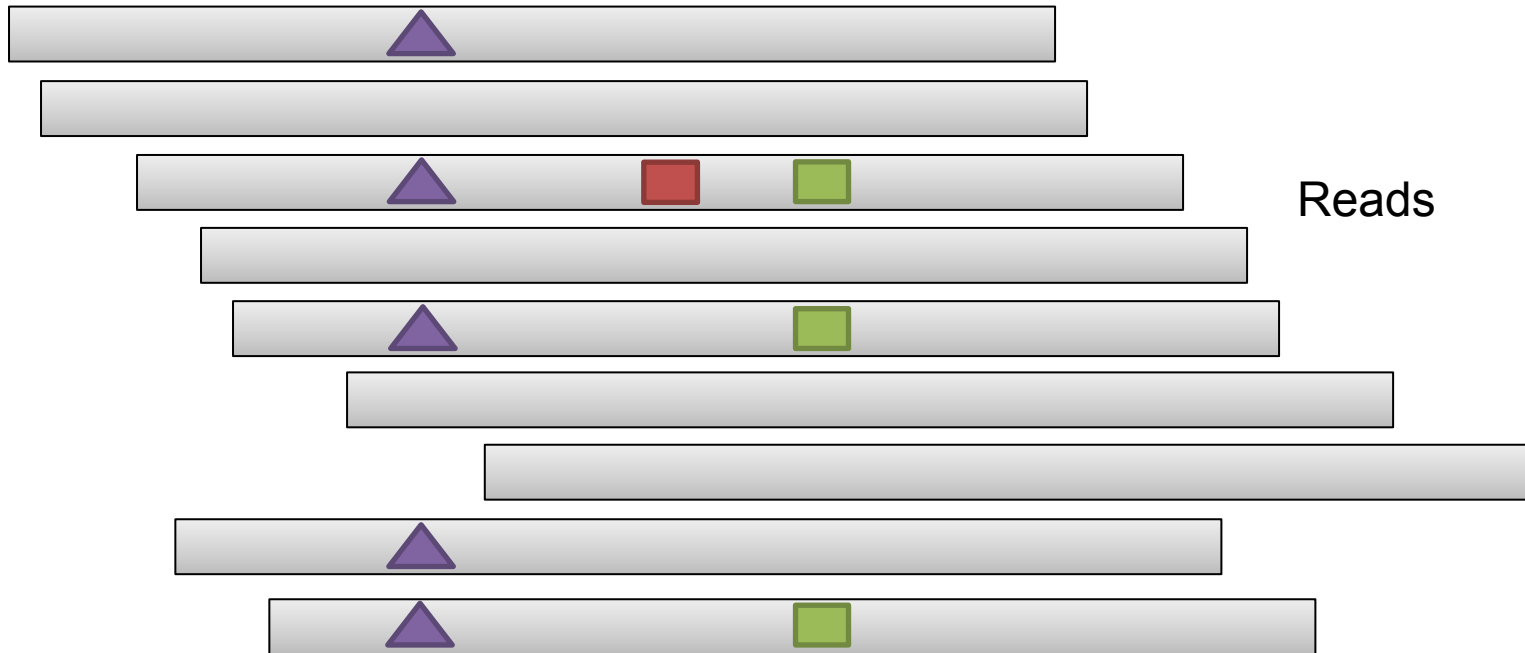


A typical workflow for alignment and variant calling.

Variant (haplotype) detection

Alignments to candidates

Reference



Reads

Variant observations

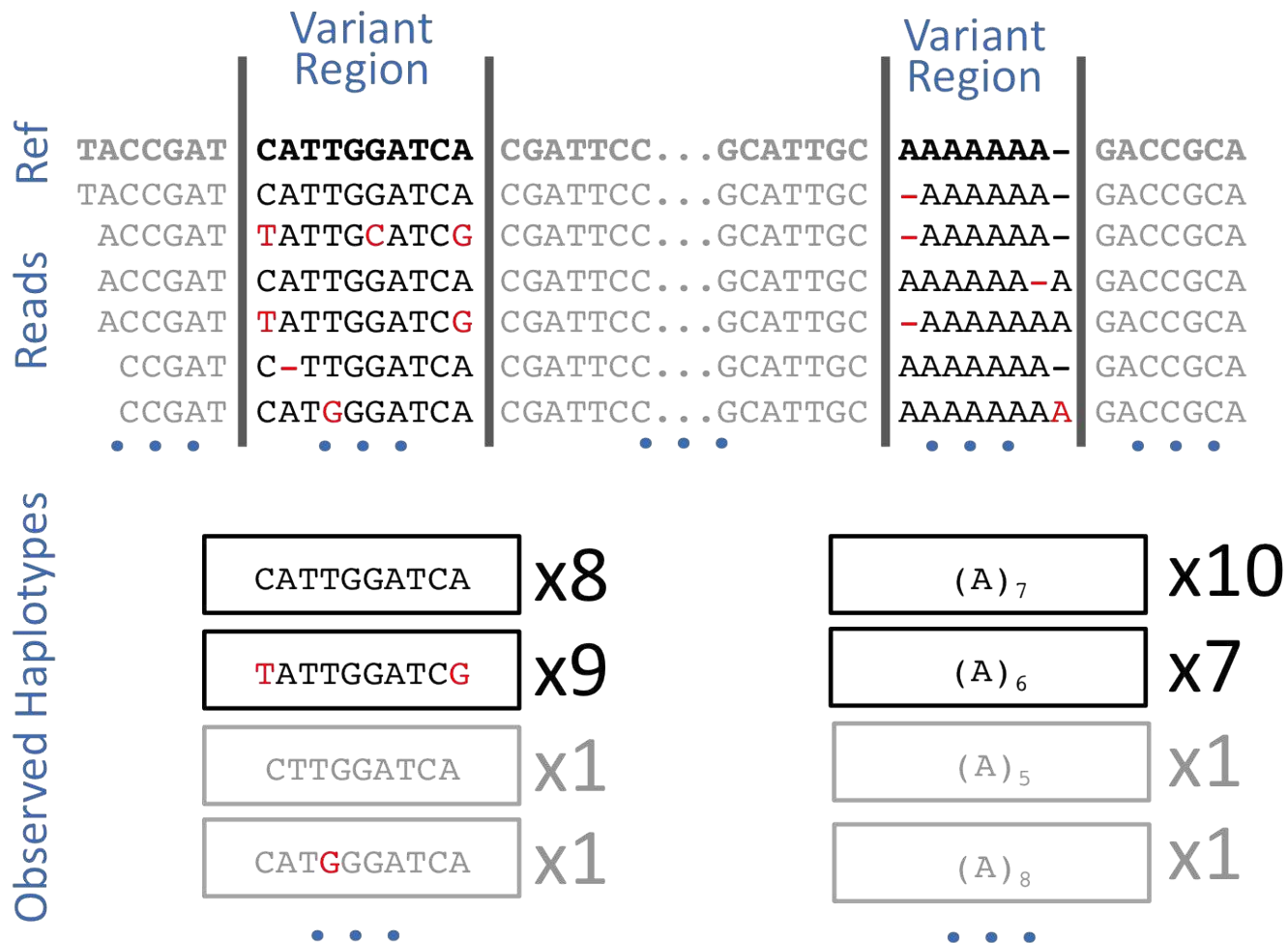
The data exposed to the caller

Reference



Direct detection of haplotypes



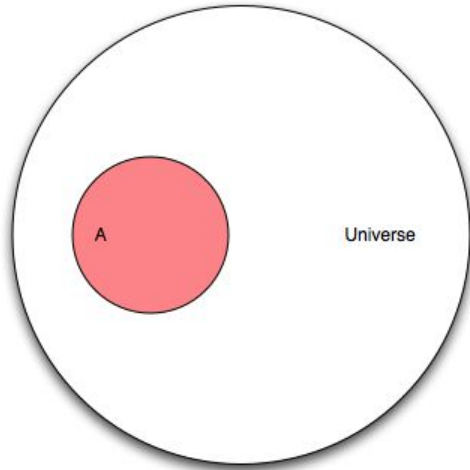


Genotyping and error detection

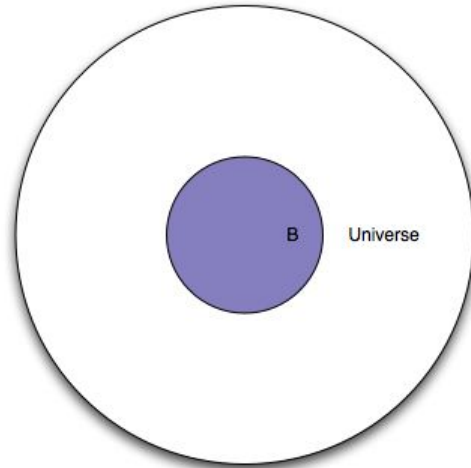
Bayesian (visual) intuition

We have a universe of individuals.

A = samples with a
variant at some locus



B = putative observations
of variant at some locus



probability(A|B)

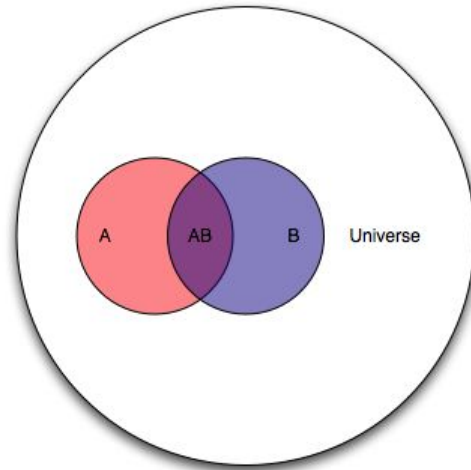
We want to estimate the probability that we have a real polymorphism "A" given "|" that we observed variants in our alignments "B".

$$P(A|B) = \frac{|AB|}{|B|}$$

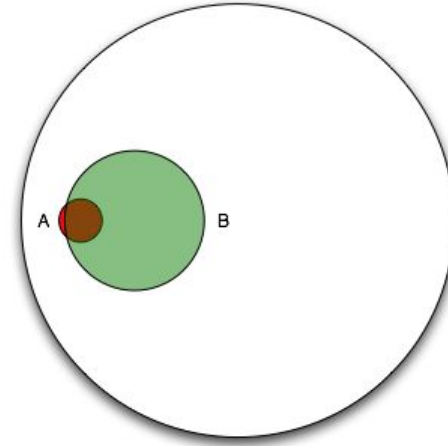
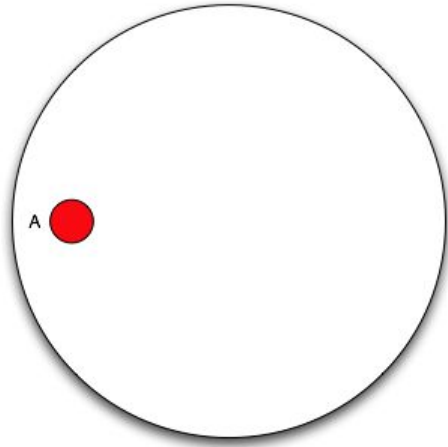
$$P(A|B) = \frac{P(AB)}{P(B)}$$

$$P(B|A) = \frac{P(AB)}{P(A)}$$

$$P(A|B) = \frac{P(B|A)P(A)}{P(B)}$$

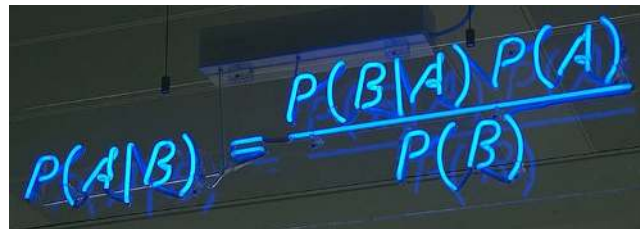


In our case it's a bit more like this...



Observations (B) provide pretty good sensitivity, but poor specificity.

The model



A photograph of a chalkboard with the formula for Bayes' theorem written in blue chalk. The formula is $P(A|B) = \frac{P(B|A) P(A)}{P(B)}$. The variables A and B are circled in blue.

- Bayesian model estimates the probability of polymorphism at a locus given input data and the population mutation rate (~pairwise heterozygosity) and assumption of “neutrality” (random mating).
- Following Bayes theorem, the probability of a specific set of genotypes over some number of samples is:
 - $\mathbf{P(G|R) = (P(R|G) P(G)) / P(R)}$
- Which in FreeBayes we extend to:
 - $\mathbf{P(G,S|R) = (P(R|G,S) P(G)P(S)) / P(R)}$
 - **G** = genotypes, **R** = reads, **S** = locus is well-characterized/mapped
 - **P(R|G,S)** is our data likelihood, **P(G)** is our prior estimate of the genotypes, **P(S)** is our prior estimate of the mappability of the locus, **P(R)** is a normalizer.

Handling non-biallelic/diploid cases

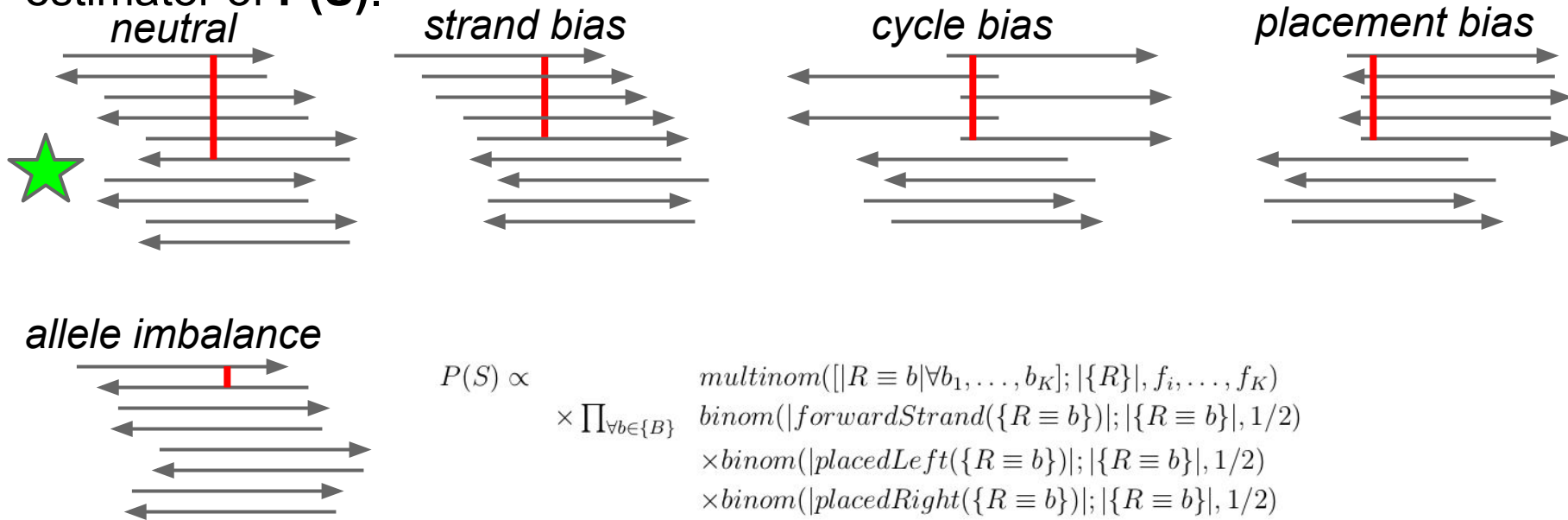
We compose our data likelihoods, **P(Reads|Genotype)** using a discrete multinomial sampling probability:

$$P(\text{reads}|\text{genotype}) = \binom{|\text{reads}|}{|\text{reads} = A|, |\text{reads} = B| \dots} \\ \times \prod_{\forall \text{alleles} \in \text{genotype}} \text{freq}(\text{allele} \in \text{genotype}) \\ \times \prod_{\forall \text{reads}} P(\text{correct}(\text{read}))$$

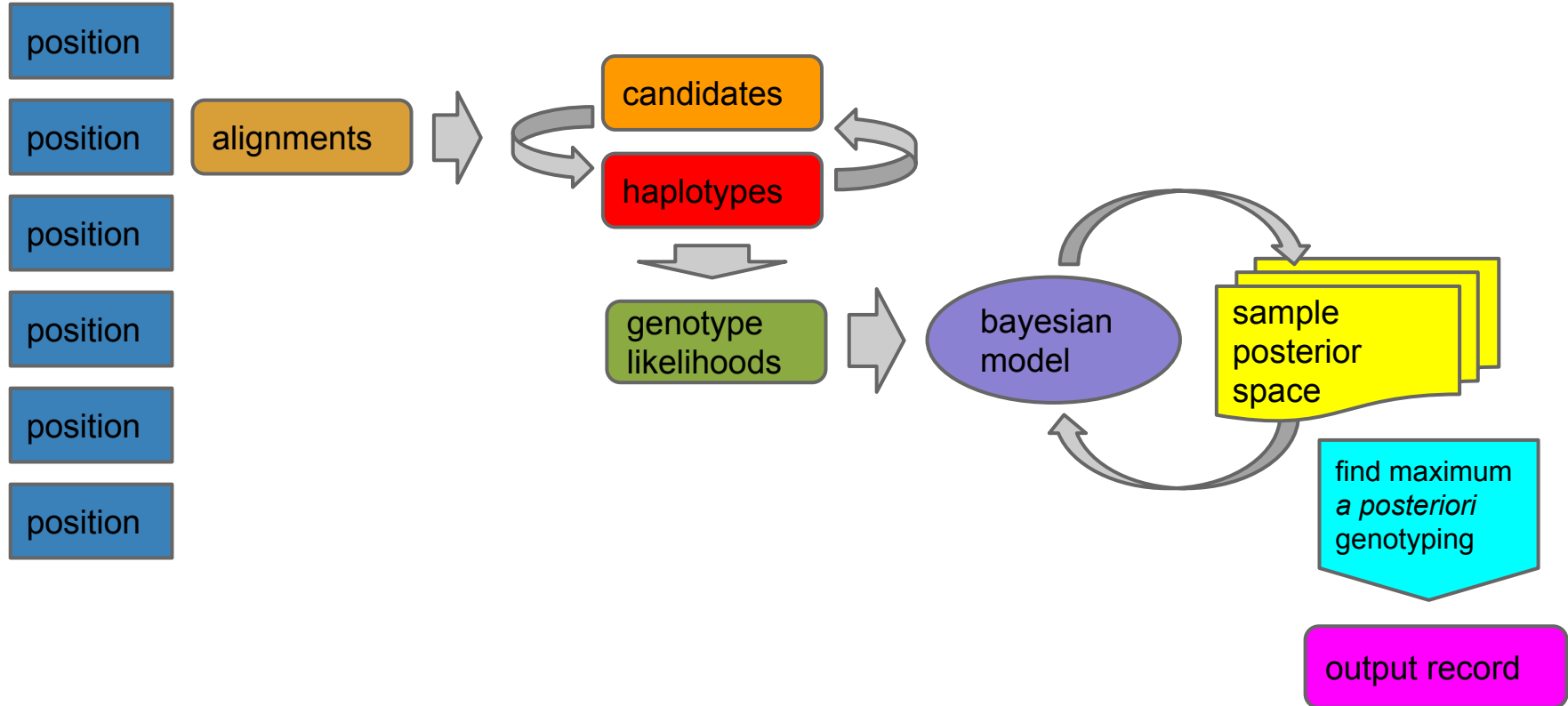
Our priors, **P(Genotypes)**, follow the Ewens Sampling Formula and the discrete sampling probability for genotypes.

Are our locus and alleles sequenceable?

In WGS, biases in the way we observe an allele (placement, position, strand, cycle, or balance in heterozygotes) are often correlated with error. We include this in our posterior $\mathbf{P}(\mathbf{G}, \mathbf{S} | \mathbf{R})$, and to do so we need an estimator of $\mathbf{P}(\mathbf{S})$.



The detection process



Variant detector lineage

PolyBayes— original Bayesian variant detector (Gabor Marth, 1999); written in perl

GigaBayes— ported to C++

BamBayes— “modern” formats (BAM)

FreeBayes— 2010-present

FreeBayes-specific developments

FreeBayes model features (~in order of introduction):

- Multiple alleles
- Indels, SNPs, MNPs, complex alleles
- Local copy number variation (e.g. sex chromosomes)
- Global copy-number variation (e.g. species-level, genome ploidy)
- Pooled detection, both discrete and continuous
- Many, many samples (>30k exome-depth samples)
- Genotyping using known alleles (hints, haplotypes, or alleles)
- Genotyping using a reference panel of genotype likelihoods
- Direct detection of haplotypes from short-read sequencing
- Haplotype-based consensus generation (clumping)
- Allele-length-specific mapping bias
- Contamination-aware genotype likelihoods

Learning to genotype

Current best practices (in humans)

Lots of people use freebayes (not in human).

FYI: The current gold standard in human genomics is DeepVariant.

It *learns* how to genotype.

We won't use it in the practicals, but you should know how it works. It could help.

We can learn to genotype

Bioinformaticians working with sequencing data can look at a visualization of alignments and make a good guess at the genotype.

right

[illegible]

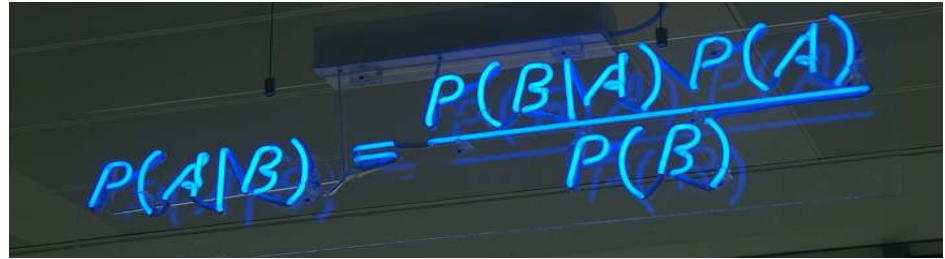
“It looks **wrong**/right”

The image displays a sequence alignment visualization on a black background with white text. The top row shows sequence positions from 141 to 241 in increments of 10. Below this, a sequence of nucleotides (G, C, A, C, A, C, A, G, C, T, A, T, C, T, C, A, A, A, C, T, T, C, T, C, A, C, A, C, T, T, T, C, C, A, A, G, C, C, C, C, T, G, A, T, C, C, C, T, T, A, T, G, A, C, T, C, T, T, G, G, C, A, G, A, T, G, C, C, C, C, A, C, C, T, T, A, T, C, T, C, T, A, C, C, A, G, A, A, C, C, T, A, A, G, G, C, C, A, A, C, T, A, G, C, A) is shown. The alignment is represented by a grid of dots and dashes. Four vertical red boxes highlight specific columns: the first box is at position 141, the second at 151, the third at 191, and the fourth at 221. The alignment shows a high degree of similarity between the sequences, with some gaps (dashes) and mismatches (dots) visible. The sequence is aligned with a reference sequence, and the alignment is shown in a way that highlights the differences between the two sequences.

Can machines learn to genotype?



Thomas Bayes

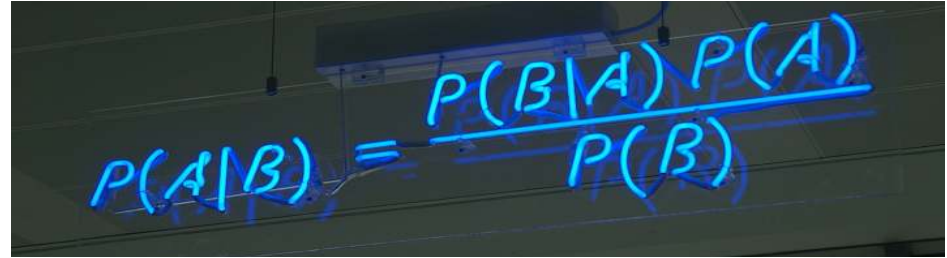
A photograph of a chalkboard with the equation for Bayes' theorem written in blue chalk. The equation is $P(A|B) = \frac{P(B|A)P(A)}{P(B)}$. The variables A and B are circled in blue.

Traditionally, we mix observations and *a priori* models using Bayesian statistics to find variants and estimate genotypes.

Can machines learn to genotype?



Marvin Minsky

A photograph of a chalkboard with the equation for Bayes' theorem written in blue chalk. The equation is
$$P(A|B) = \frac{P(B|A)P(A)}{P(B)}$$

$$P(A|B) = \frac{P(B|A)P(A)}{P(B)}$$

But instead of building our model and prior from first principles, we could learn it (with machines).

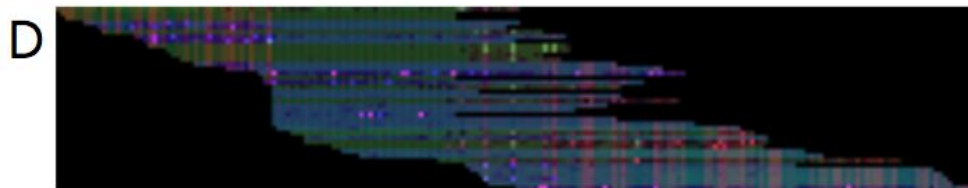
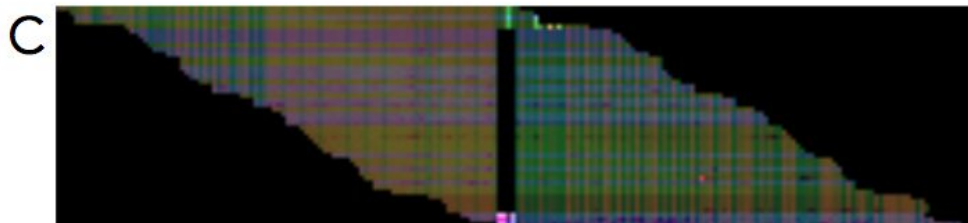
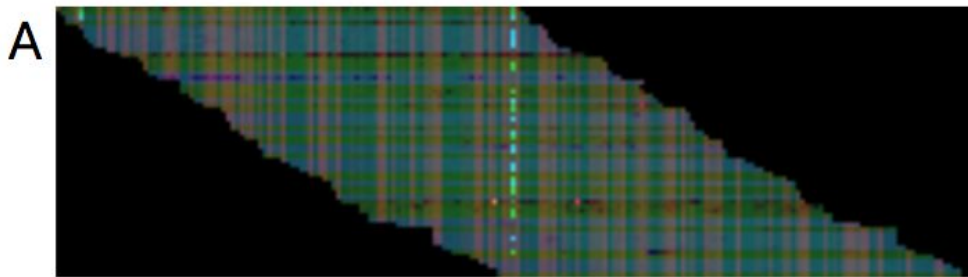
DeepVariant

Idea: you can leverage convolutional neural networks designed for images to *learn to genotype*.

<https://doi.org/10.1038/nbt.4235>

DeepVariant inputs

Idea: convert alignments to the reference into read pileups. Annotate various channels with useful things (like quality, read base, reference base, etc.)

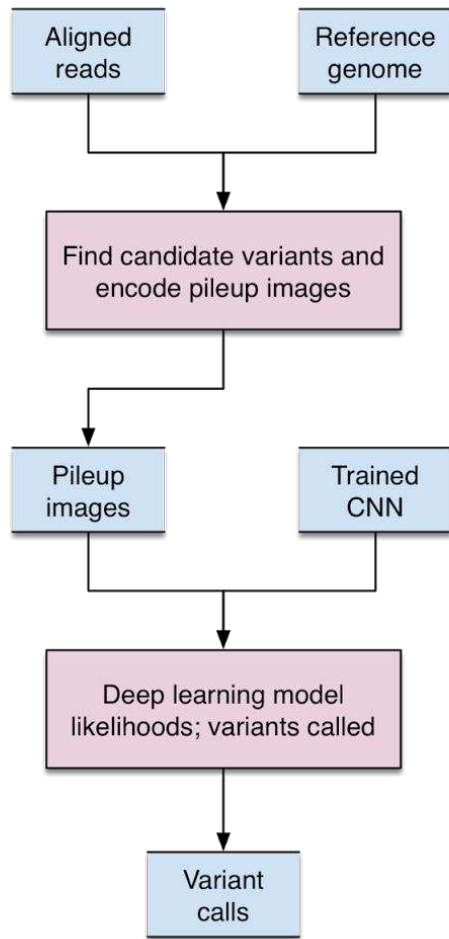


Channels are shown in greyscale below in the following order:

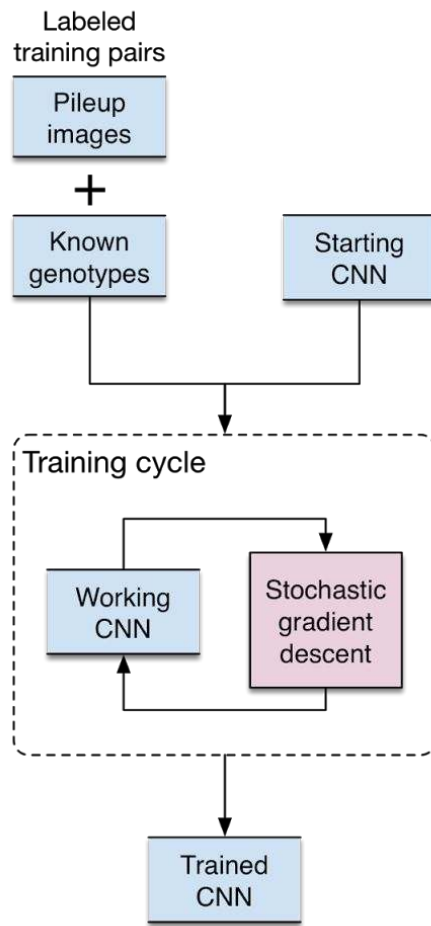
1. Read base: different intensities represent A, C, G, and T.
2. Base quality: set by the sequencing machine. White is higher quality.
3. Mapping quality: set by the aligner. White is higher quality.
4. Strand of alignment: Black is forward; white is reverse.
5. Read supports variant: White means the read supports the given alternate allele, grey means it does not.
6. Base differs from ref: White means the base is different from the reference, dark grey means the base matches the reference.



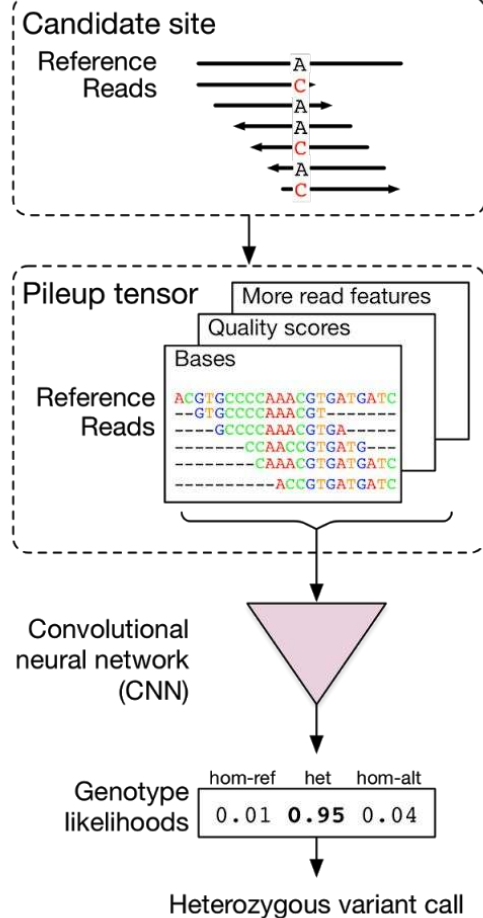
DeepVariant



DeepVariant CNN training



Pileup image evaluation



Words of caution

- Neural networks are universal approximators.
- But, they're only guaranteed to model any pattern over the domain in which they've been trained.
- DeepVariant may work well (it wins all variant calling competitions), but it's worth comparing it to other methods in the context of non-human genomes.

Practicals

Practical

Walkthrough:

<https://github.com/ekg/alignment-and-variant-calling-tutorial/tree/evomics2024>

Goal is to get our hand dirty with a variant calling workflow, and maybe to dig into interesting edge cases that arise.

Notes

- The workflow is pretty complete. You can copy-paste if you want, but please look at what you're ingesting and producing at each step.
- Data and time-consuming indexing results are in `~/workshop_materials/variant_calling`
- If something is taking too long or you've got stuck, there is also a `.results` directory that contains most outputs.
- At the end we've got some open-ended mini-projects. Explore them!

