

Systematic Substitution Errors Found in...



mlelivel

Systematic Substitution Errors Found in MiSeq (R) Data Jan 31, 2012 9:32 AM

Systematic Homopolymer Errors in MiSeq® Sequencer Data

Mike Lelivelt & Yutao Fu, Ion Torrent by Life Technologies

Semiconductor sequencing by Ion Torrent® offers post-light DNA sequencing using unmodified nucleotides and naturally occurring DNA polymerase. This biochemistry more closely resembles the DNA replication process that has evolved over billions of years. Illumina's MiSeq® DNA sequencer relies on nucleotides that contain modified fluorescent molecules that must be both excited (via expensive light sources) for detecting the base and then completely removed from the growing polymer chain so nucleotide incorporation can continue. This leads to the following question: does the partial or incomplete removal of these fluorescent molecules impact the accuracy of the MiSeq® sequencer?

Observation: The MiSeq® sequencer tends to mis-call bases adjacent to a homopolymer region, especially stretches of Cs and Gs, thus falsely extending the homopolymer length and causing strand-specific substitution errors.

To illustrate the phenomenon, a small region of the *E. coli* genome was loaded into a standard genome browser, such as the Broad's IGV, using both public MiSeq® and Ion Torrent® *E. coli* DH10B datasets^[1,2]. Figure 1 (A&B) shows a particular region of DH10B genome in which substitution errors are concentrated in MiSeq® reads, but not Ion Torrent® reads (such genomic regions aren't difficult to find in DH10B or K12). A careful review of this data reveals an interesting pattern of errors tending to replicate the previous (5' upstream) base calls. Following a homopolymer stretch in MiSeq®, the base that follows this stretch tends to be erroneous, and those errors tend to retain the preceding

base from the homopolymer stretch. These substitution errors often fall to the last base of a homopolymer region - based on the direction of the read. For example, in a stretch of three G bases, the fourth base is often erroneously called a G. This strand-specific pattern is wide spread, and explains 49.9% and 51.8% of MiSeq® substitution errors overall in DH10B and K12, respectively. This dominant error profile that can be found so frequently next to homopolymer regions suggests a clear systematic bias within MiSeq® data.

Measuring the trend across the whole genome - conditional probability

In order to measure this trend across a larger genomic region, the conditional probability $p(C|Cx)$ was measured across homopolymer regions across all bases sequenced. This conditional probability is the chance of observing another C base following runs of x Cs, and similarly $p(G|Gx)$ for G. In other words, it is simply likelihood that the next base in a sequence will be called the same as the previous base. A baseline conditional probability is calculated based on the naturally occurring genomic sequence. The closer the conditional probability of each sequencing technology is to the genome background value, the less bias exists. The conditional probabilities were calculated for *E. coli* K12 and DH10B genomes for both MiSeq® and Ion Torrent® publicly available data sets. Bases not aligned to the genome were excluded as they are disregarded from the perspective of calling variants. For both K12 and DH10B, MiSeq® datasets showed much larger conditional probability deviations from genomic background than the matching Ion Torrent® datasets, as shown in Figure 2A & 2B. The deviation grows dramatically as homopolymer length increases out to 5 bases. Similar but less severe MiSeq® bias was also observed for bases A and T.

Sequencing *E. coli* is an excellent control substrate to test the accuracy of DNA sequencing. However, researchers often want to see sequencing performance measured on human DNA. The conditional probabilities of both C and G bases were measured from publicly available human amplicon data sets from both MiSeq®^[3] and Ion Torrent PGM® sequencers. The same bias of overcalling G/C bias was also observed in the MiSeq® human amplicon dataset (Figure 2C). The trend was especially noticeable in long stretches of both G and C bases.

Hypothesizing why such a trend exists?

One hypothesis that fits nicely with the above observations: the fluorescent dyes used by MiSeq® sequencing are subject to incomplete cleavage and may accumulate across a homopolymer region. If fluorescent dyes are retained from the previous flow and then excited in the subsequent flow, a wrong base call is more likely to be produced. Such a symptom may be unique to reversible terminator based sequencing chemistry. Independent studies on Illumina sequencing data already suggested filtering out “variants flanked by homopolymers of length greater than 3 or surrounded by greater than 6 identical bases”^[4].

Conclusion

Light-based DNA sequencing from MiSeq® Sequencer is only as good as the ability for the technology to uniformly remove the fluorescent moiety from the growing polymer at each flow cycle. The data available from multiple public releases of MiSeq® data suggest that MiSeq® data contains systematic errors associated with homopolymer regions and the impact of these errors is larger in the MiSeq® platform relative to the Ion Torrent PGM® platform especially in G/C homopolymer regions.

Figure 1. IGV screenshots of MiSeq® (panel A - top) and Ion Torrent® (panel B - bottom) datasets for a genomic region of *E. coli* DH10B

Systematic Substitution Errors Found in...

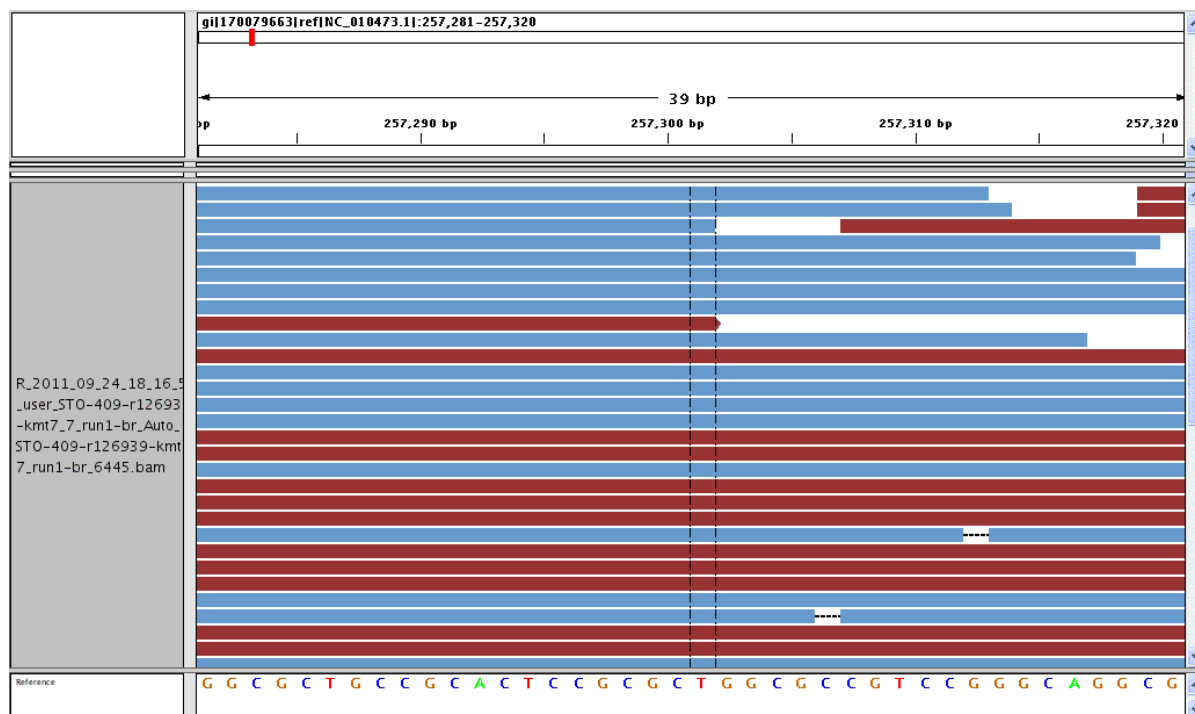
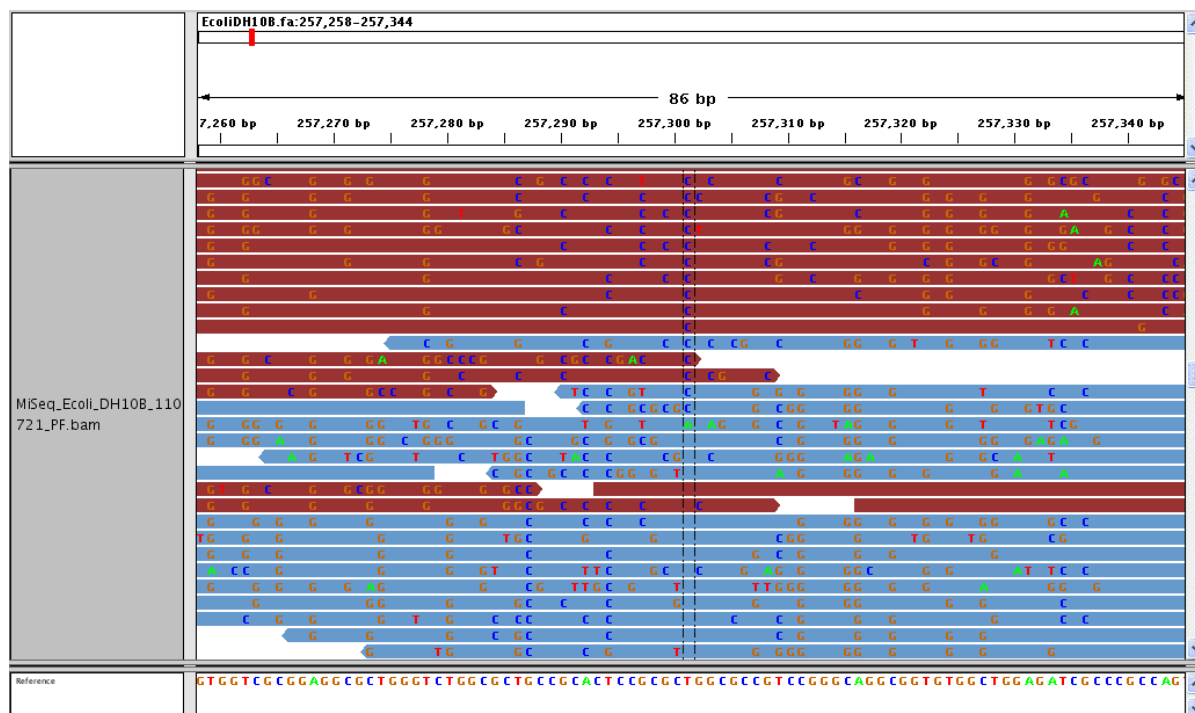
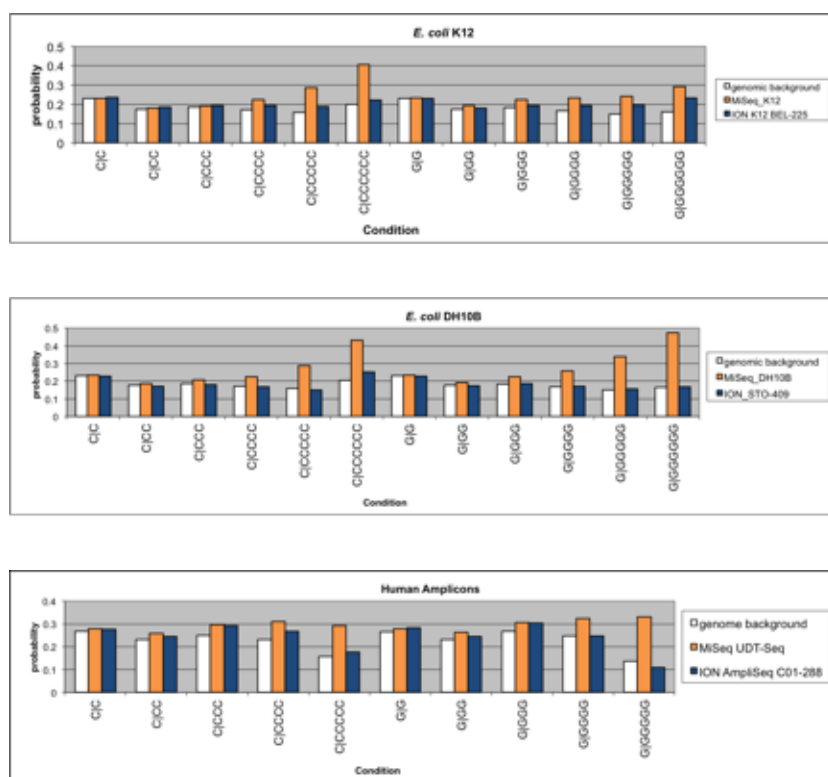


Figure 2. Conditional probabilities $p(C|Cx)$ and $p(G|Gx)$ for increasing homopolymer length x . Panels A (top) for *E. coli* K12, B (middle) for *E. coli* DH10B, and C (bottom) for Human Amplicons. For panel C, only the genomic background distribution for Ion targeted genome regions is shown. MiSeq background from UDT-Seq was calculated independently, as the composition of the amplicons between the panels is different. The difference between the background probabilities was not significantly different.



Methods:

BAM files for *E. coli* data were downloaded from Ion Community (http://lifetech-it.hosted.jivesoftware.com/community/torrent_dev, except for a K12 run from an R&D Torrent Server) and Illumina web pages (http://www.illumina.com/systems/miseq/scientific_data.ilmn). MiSeq® human amplicon reads were downloaded from NCBI SRA

page for run SRR385941, converted to fastq format, and mapped to target regions to produce a BAM file. Both BWA and a C program in Torrent Suite, TMAP, were used for mapping, and produced similar observations downstream. The parameters used for TMAP were:

```
tmap mapall -n10 -f targets.fasta -r SRR385941.fastq map1 map2 map3|samtools view -bS -  
| samtools sort - results
```

, and for BWA:

```
bwa aln -q 15 targets.fasta SRR385941.fastq > SRR385941.sai
```

```
bwa samse targets.fasta SRR385941.sai SRR385941.fastq > results.bam
```

The bam files were first parsed first using a C++ program alignStats, also available as a binary in Torrent Suite, with the following parameters:

```
alignStats -i <bam_file> -p 1
```

The resulting Default.sam.parsed files were used for conditional probabilities.

A PERL script 'countpattern.pl' (appendix 1) was used to search aligned portions of sequencing reads recorded in Default.sam.parsed files for HxM patterns, in which Hx is a homopolymer with $x(x \geq 1)$ base of H (H=A,T,C or G), and M is a monomer base. The counts were tallied in Excel and conditional probabilities of homopolymer extension $p(H_{x+1}|H_x)$ were inferred as the ratio between counts of $\{H_xM|H=M\}$ and $\{H_xM\}$.

Appendix:

1. countpattern.pl:

```
#!/usr/bin/perl

$pattern = $ARGV[1]?$ARGV[1]:'CG';

$patternlen = $ARGV[2]?$ARGV[2]:length($pattern);

open in, "$ARGV[0]";

while (<in>)

{if (/>(.*)/) {push @s,$1;$s=$1;}

else {s/ //g; $s{$s}.=uc($_)}

}

close in;

for (@s)

{$seq = uc($s{$_});

pos $seq=0;

$count=0;

while ($seq=~/$pattern/g)

{$count++;

pos($seq) = pos($seq)-$patternlen+1;

}

$len = length($seq);

$len++ if $len==$patternlen-1;

print "$_\\t",int($count/($len-$patternlen+1)*10000+0.5)/100,"\\t$count\\n";
```

}

__END__

#Example usage:

```
for a in C CC CCC CCCC CCCCC CCCCCC G GG GGG GGGG GGGGG GGGGGG;  
do for b in A C G T;do echo -ne "$a\t$b\t" >> output.txt; countpattern.pl  
sample.fasta $a$b|cut -f3 >> output.txt; done; done;
```

References:

1. http://www.illumina.com/systems/miseq/scientific_data.ilmn
2. http://lifetech-it.hosted.jivesoftware.com/community/torrent_dev (except a K12 dataset which is not yet released)
3. Olivier Harismendy, Richard B Schwab, Lei Bao, Jeff Olson, Sophie Rozenzhak, Steve K Kotsopoulos, Stephanie Pond, Brian Crain, Mark S Chee, Karen Messer, Darren R Link and Kelly A Frazer. Detection of low prevalence somatic mutations in solid tumors with ultra-deep targeted sequencing.

Genome Biology 2011, 12:R124

4. Larson DE, Harris CC, Chen K, Koboldt DC, Abbott TE, Dooling DJ, Ley TJ, Mardis ER, Wilson RK, Ding L. SomaticSniper: Identification of Somatic Point Mutations in Whole Genome Sequencing Data. Bioinformatics (2011) epub Dec 6, doi: 10.1093/bioinformatics/btr665

Tags: miseq_homopolymer



flxlex

Systematic Substitution Errors Found in MiSeq (R) Data Jan 30, 2012 11:27 AM

Interesting results...

As DH10B is a substrain of E. coli K12, perhaps what you call 'K12' is E. coli K12 substrain MG1655? Could you please (please) make the run for this strain available? This would finally make a comparison with 454 E. coli data (always generated from MG1655) possible...



markus_g

Systematic Substitution Errors Found in MiSeq (R) Data Jan 31, 2012 9:51 AM

Hi,

Is this maybe related to this here?

Nakamura K, Oshima T, Morimoto T, et al. Sequence-specific error profile of Illumina sequencers. Nucleic acids research. 2011;39(13):e90.

The MIRA assembler could already handle these errors long before that paper was published.



lek2k

Systematic Substitution Errors Found in MiSeq (R) Data Jan 31, 2012 3:33 PM

Hey thanks for the NAR reference. Interesting read..



lek2k

Systematic Substitution Errors Found in MiSeq (R) Data Jan 31, 2012 3:32 PM

A similar but independent analysis performed on the MiSeq Bacillus cereus genome.

<http://biolektures.wordpress.com/2012/02/01/are-miseq-miscalls-influenced-by-preceding-homopolymers/>