

Helicobacter pylori Genome Project (HpGP)

Difei Wang, Ph.D.

**Bioinformatics Manager
Research Analysis Support Group
Cancer Genomics Research Laboratory
Frederick National Laboratory for Cancer Research
Leidos Biomedical Research, Inc.**

2024-02-21

Supporting the Division of Cancer Epidemiology and Genetics



H. pylori is a carcinogen



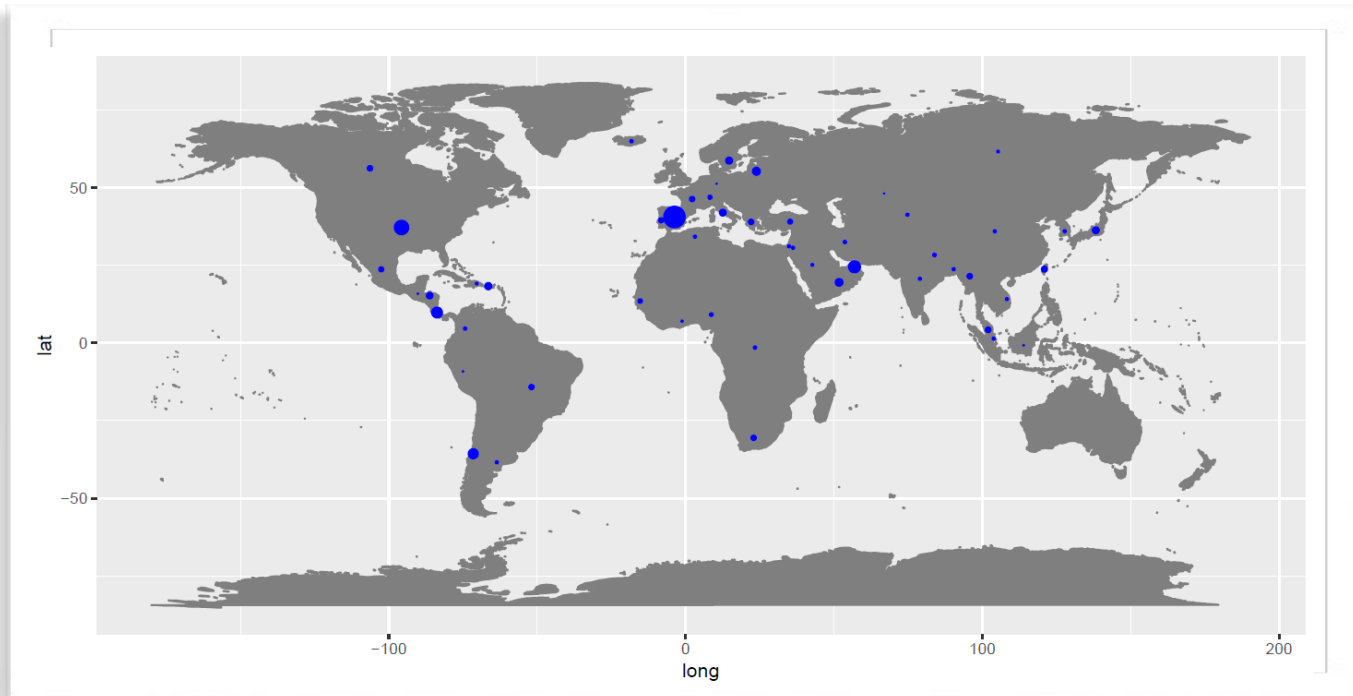
- *H. pylori* is a gram-negative bacterium and has been linked with stomach cancer. About 2 out of 3 adults worldwide are infected with *H. pylori*.
- Still limited recognized virulence factors (e.g. *cagA* and *vacA*, highly prevalent in high-risk areas, e.g. East Asia)
- Genome first sequenced in 1997. Genome size is about 1.6 Mb. About 1,500 coding genes.
- A few hundred complete genomes in GenBank in 2016. Most of them from RSII or Illumina short-reads.
 - Still a limited number of genomes per geographic site
- Genome-wide base modifications first profiled in 2014 (“Methylome”)
- Newly available technology
 - Single molecule, real-time (SMRT) sequencing: PacBio RSII, Sequel and Sequel II and Nanopore platforms.

HpGP was launched in 2017.

Co-PIs: Dr. Constanza Camargo
Dr. Charles Rabkin

<https://www.cancer.gov/about-cancer/causes-prevention/risk/infectious-agents/h-pylori-fact-sheet#r23>

~ 51 HpGP Collaborating Centers



Each Center to Contribute:

Gastric Cancer (GC)
Adv. Intestinal Metaplasia (IM)
Non-atrophic Gastritis (NAG) and others

Total more than 1033 samples collected and sequenced so far.



HpGP Objectives

- To map the global population structure of > 1000 *de novo* assembled complete *H. pylori* genomes
- To characterize the spectrum of genomic and epigenomic variations of benign strains of *H. pylori*
- To identify molecular features that may contribute to pathologic effects
- To establish a resource repository of multidimensional data and well-characterized strains for utilization by the scientific community

Co-PIs: Dr. Constanza Camargo
Dr. Charles Rabkin

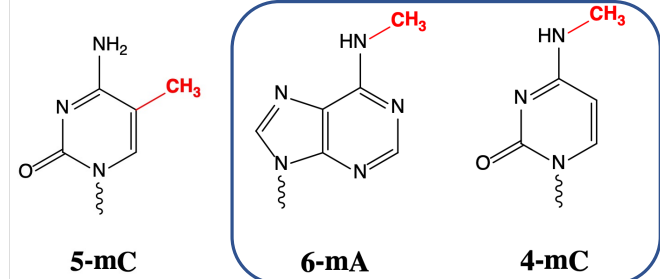
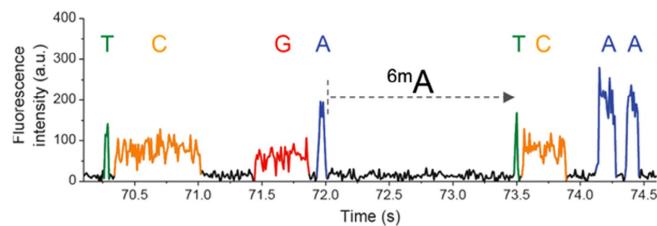
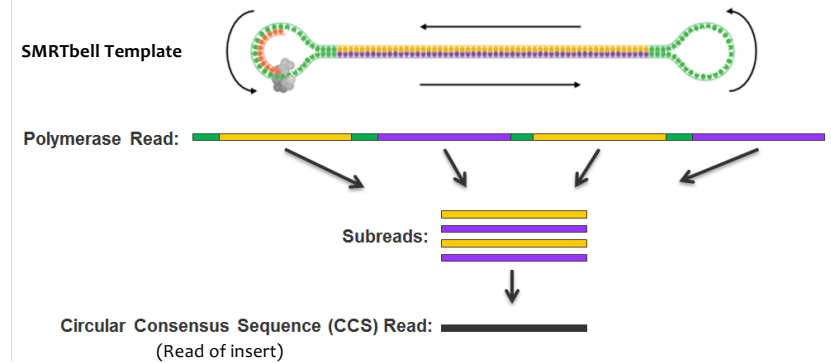
Single Molecule, Real-Time (SMRT) Sequencing



RSII



Sequel/Sequel II



<https://www.pacb.com/>

A Pilot Study in 2017

- Mexico: 43 [23 Cancer, 17 Gastritis, 3 Others]
- Honduras: 27 [8 Cancer, 9 NAG, 10 IM]
- **Colombia:** 26 [2 Cancer, 12 IM, 12 NAG]
- Lativa: 21 [1 Cancer, 10 NAG, 10 IM]
- Taiwan: 20 [10 Cancer, 10 NAG]

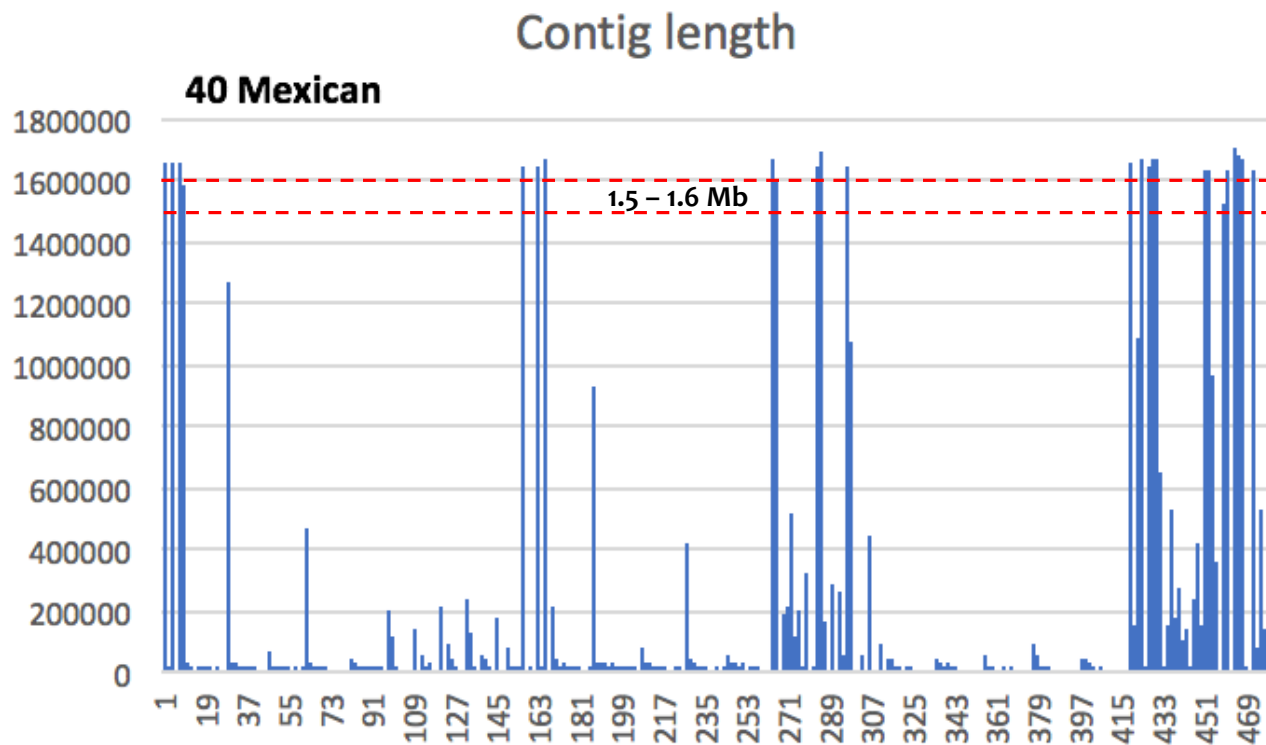
- Greece: 8
- Peru: 10

~ 155 *H. pylori* genome sequences were sequenced on PacBio RSII, de novo assembled and delivered. (Frederick)



In order to know how well they were assembled, we need to do some quality checking and align the assembled genome sequences to the know sequences.

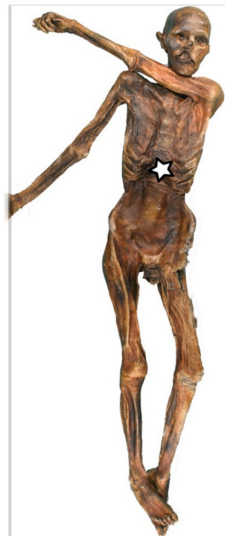
Number and Length of Contigs



~ 40%

With more
than 1 contig
and each of
them much
less than
1.5Mb

Known *H. pylori* Genomes



ANCIENT MICROBIOME

The 5300-year-old *Helicobacter pylori* genome of the Iceman

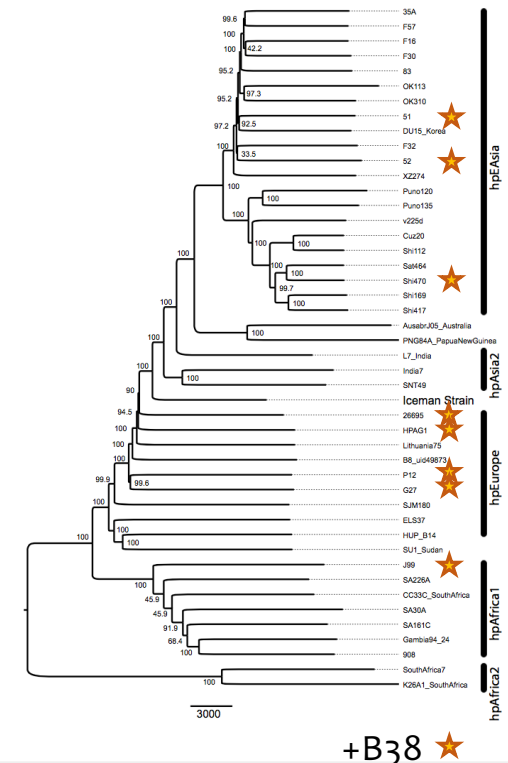
Frank Maixner,^{1,*†} Ben Krause-Kyora,^{2†} Dmitriy Turaev,^{3†} Alexander Herbig,^{4,5†} Michael R. Hoopmann,⁶ Janice L. Hallows,⁶ Ulrike Kusebauch,⁶ Eduard Egarter Vigl,⁷ Peter Malferttheiner,⁸ Francis Megraud,⁹ Niall O'Sullivan,¹ Giovanna Cipollini,¹ Valentina Coia,¹ Marco Samadelli,¹ Lars Engstrand,¹⁰ Bodo Linz,¹¹ Robert L. Moritz,⁶ Rudolf Grimm,¹² Johannes Krause,^{4,5,†} Almut Nebel,^{2,†} Yoshan Moodley,^{13,14,†} Thomas Rattei,^{3,†} Albert Zink^{1,*,†}

The stomach bacterium *Helicobacter pylori* is one of the most prevalent human pathogens. It has dispersed globally with its human host, resulting in a distinct phylogeographic pattern that can be used to reconstruct both recent and ancient human migrations. The extant European population of *H. pylori* is known to be a hybrid between Asian and African bacteria, but there exist different hypotheses about when and where the hybridization took place, reflecting the complex demographic history of Europeans. Here, we present a 5300-year-old *H. pylori* genome from a European Copper Age glacier mummy. The "Iceman" *H. pylori* is a nearly pure representative of the bacterial population of Asian origin that existed in Europe before hybridization, suggesting that the African population arrived in Europe within the past few thousand years.

Oetzi found in Italy in 1991

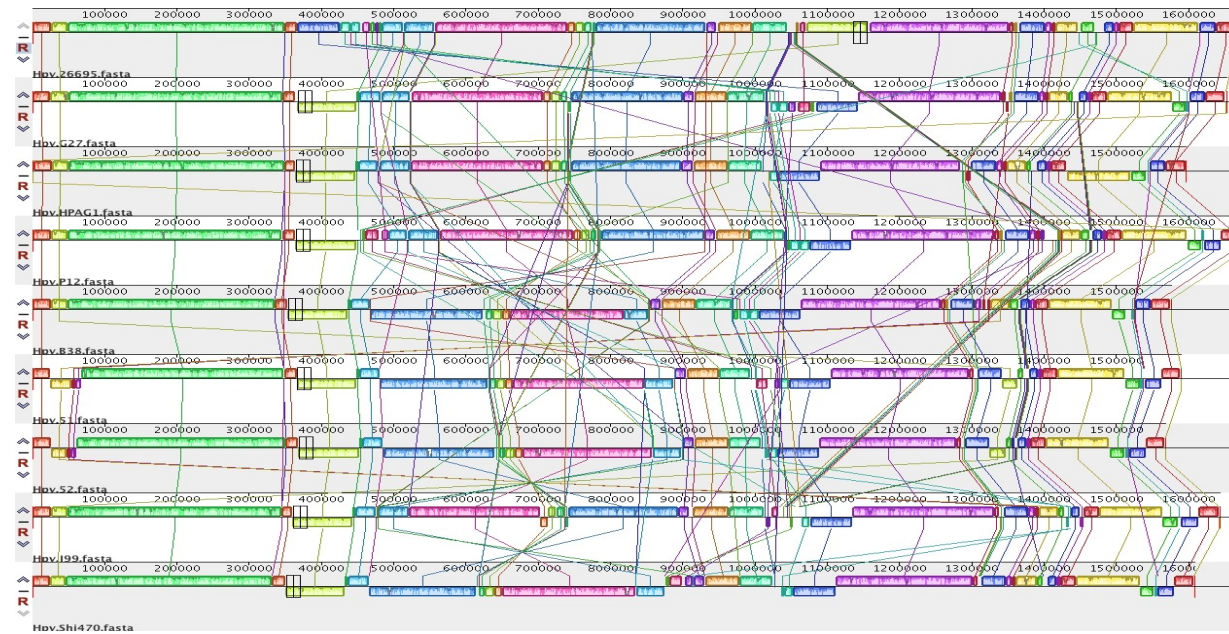
hpAsia2 (Indian & Europe), cagA+, vacA+

Selected 9 complete *H. pylori* genomes from NCBI for alignment.



Genome Alignment of 9 Published *H. pylori* Genomes in GenBank

hpEurope
 26695
 G27
 HPAG1
 P12
 B38
hspEAsia
 51
 52
hspWAfrica
 J99
hspAmerInd
 Shi470



Genome Alignment of 12 Colombian (NAG) *H. pylori* Genomes to HP26695

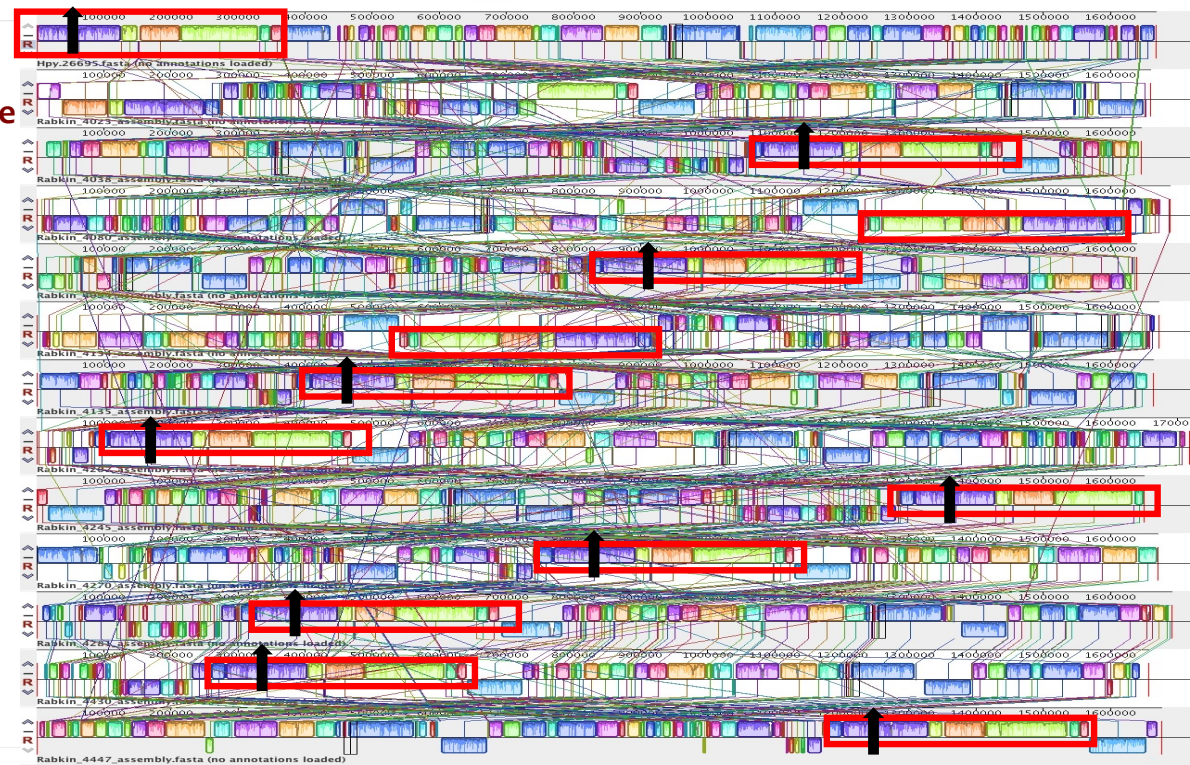
hpEurope 26695

Colombia_NoChange

Rabkin_4023

....

Rabkin_4447



Genome Alignment of 12 Colombian (Progression) *H. pylori* Genomes to HP26695

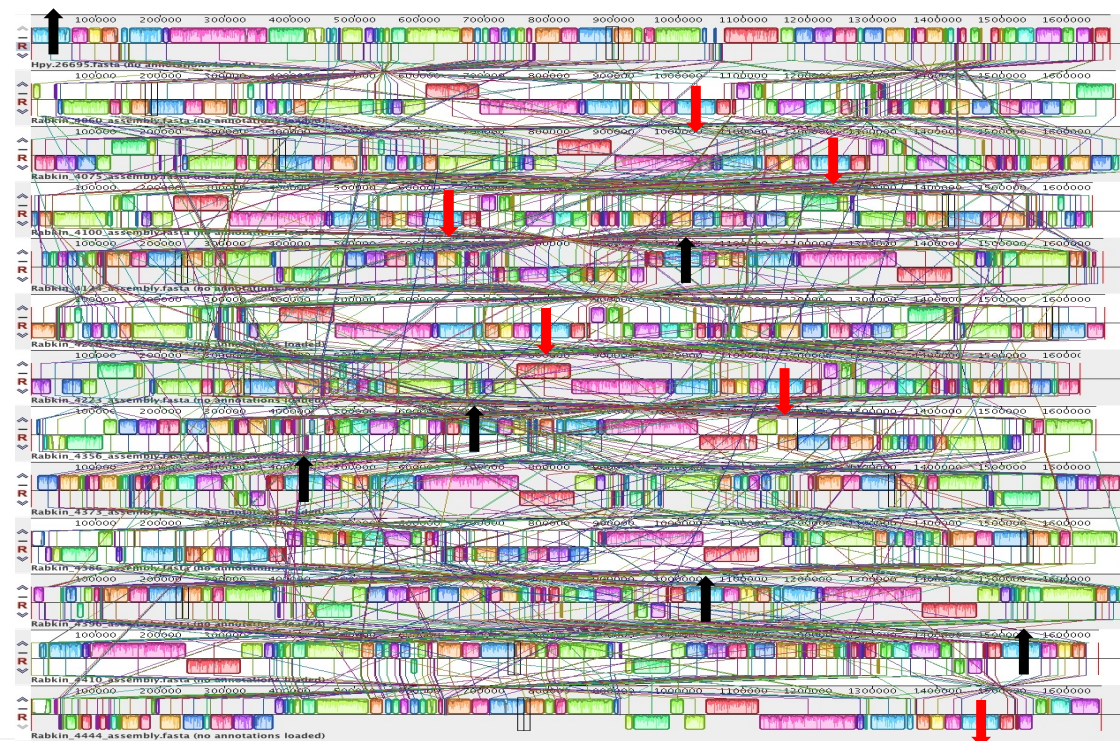
hpEurope 26695

Colombia_Progression

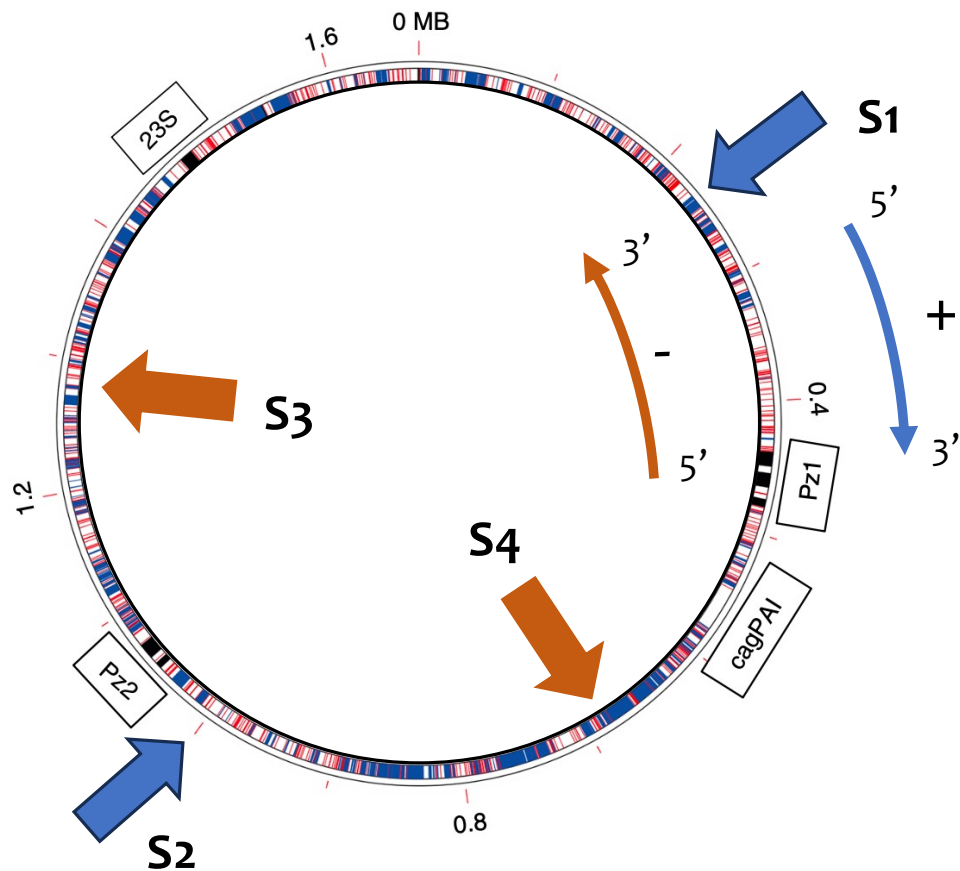
Rabkin_4060

....

Rabkin_4444



What Happened?



H. pylori genome is circular.

De novo assembling gives you the assembled sequence with inconsistent starting sites/points.

- ❖ For example, the assembled sequence can start at S1, S2, S3 or S4.

So re-ordering the start point of the genome sequence is needed!

How to do it in a proper way?

- ❖ First gene selection etc.

Re-ordering the Genome

1. Find the nusB gene in the assembled genome

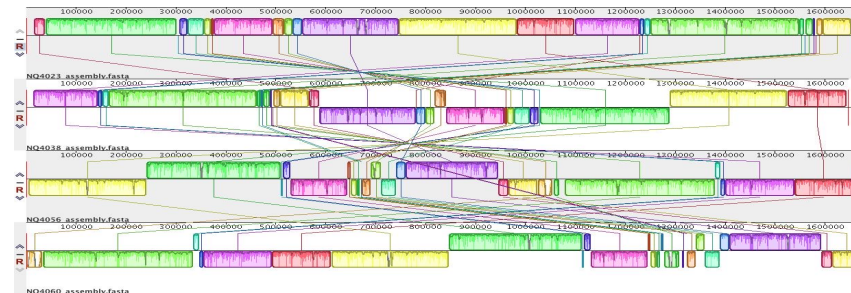


2. Set it as the first gene in genome

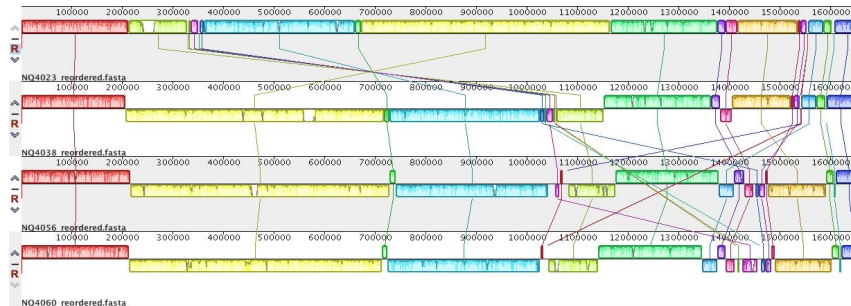


3. Make the cut and join two pieces of sequences together

Before re-order



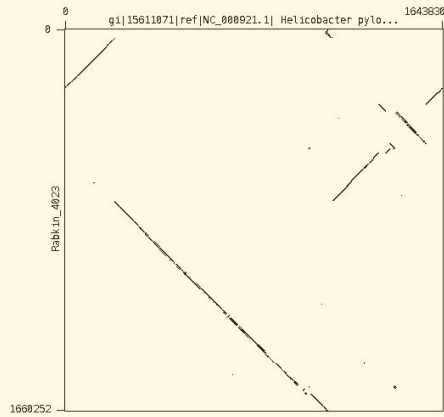
After re-order



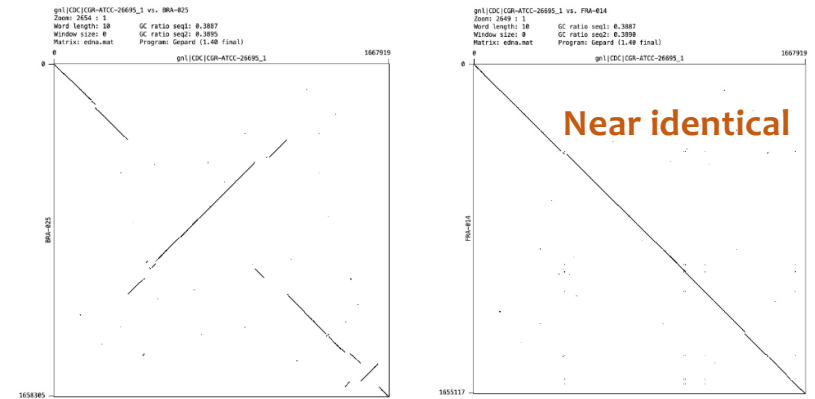
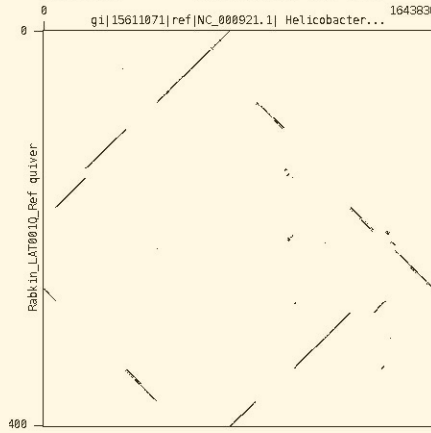
Dr. Kasia Thorell
Karolinska Institute

DOT plot

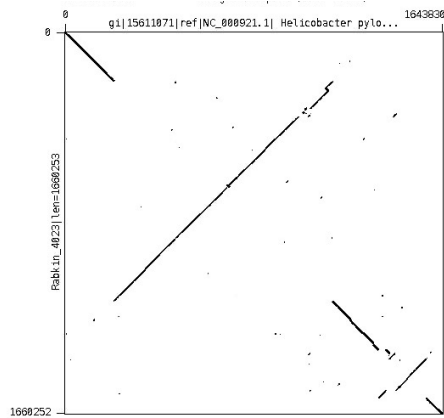
HpyJ99 vs. COL4023



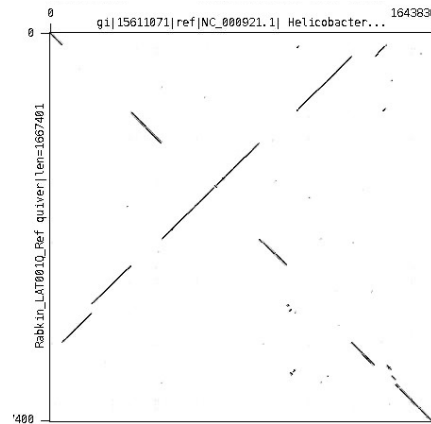
HpyJ99 vs. LAT001



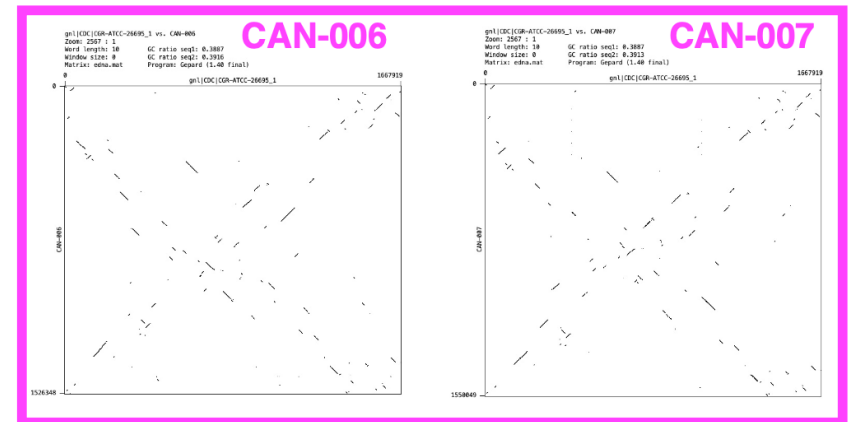
HpyJ99 vs. COL4023 fixed



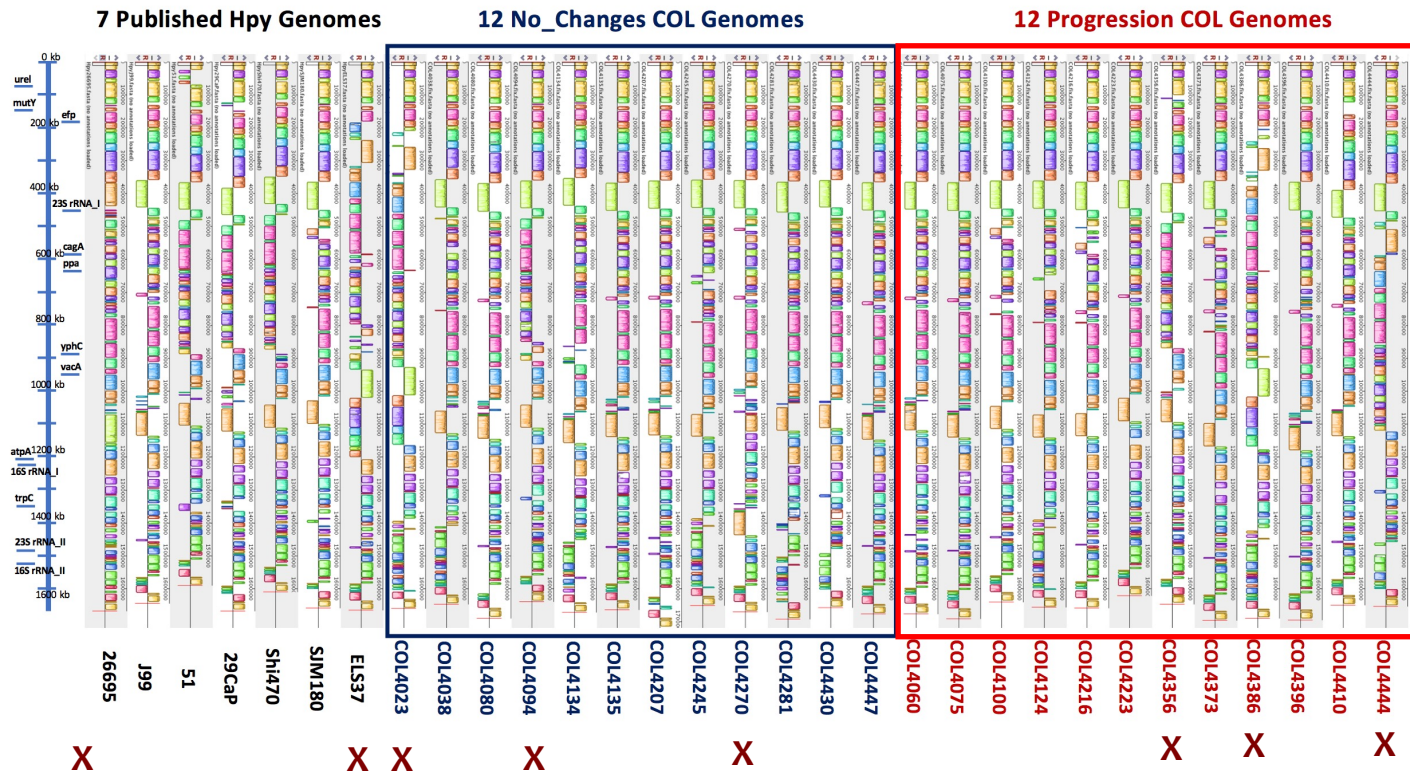
HpyJ99 vs. LAT001 fixed



Outliers



7 Published and 24 COL Genomes Alignment



Mauve genome alignment



Other Major Issues and Challenges

- In 2017, most groups used Illumina technology to do *de novo* assembly or align reads to HP26695.

Quite often, they got incomplete assemblies.

- How do we check the sequencing quality?
- How do we evaluate these *de novo* assemblies?
- How do we find the unintentional duplicates in the dataset?
- How to evaluate plasmids?

Prof. Ichizo Kobayashi, University of Tokyo

Dr. Richard Roberts, New England Biolabs

Internal investigators

Wen, Kedest, Kristie
Josh Cherry(NCBI), John Dekker(NIAID)

Other Major Issues and Challenges



- In 2017, most groups used Illumina technology to do *de novo* assembly or align reads to HP26695.

Quite often, they got incomplete assemblies.

- Comparison of the Illumina vs PacBio results

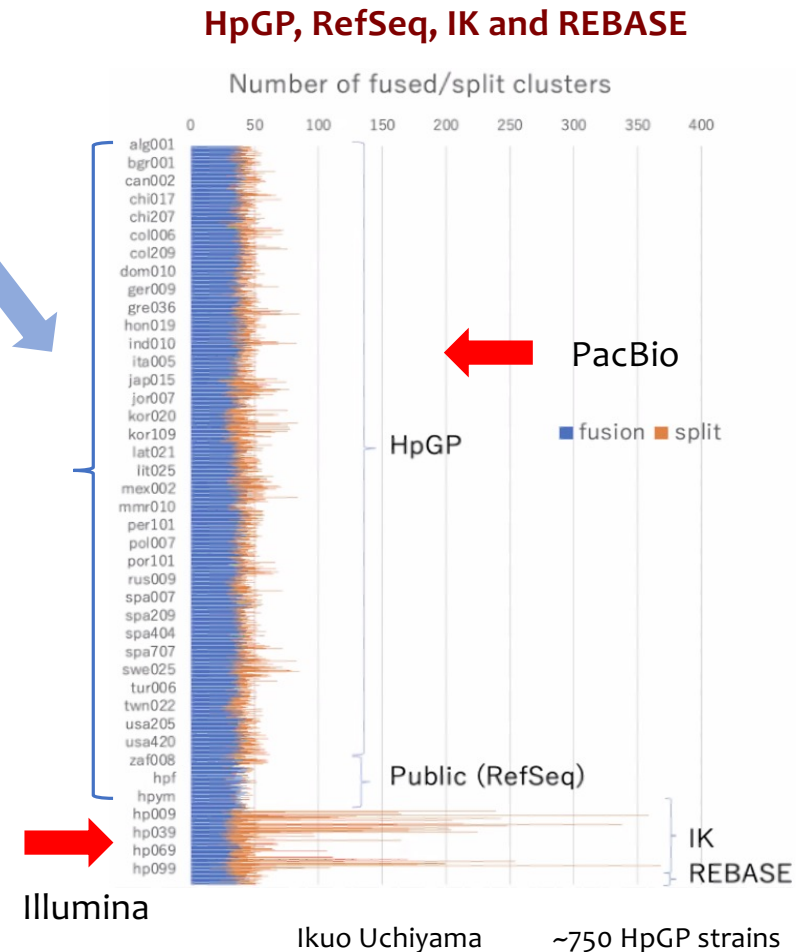
- How do we check the sequencing quality?
- How do we evaluate these *de novo* assemblies?
- How do we find the unintentional duplicates in the dataset?
- How to evaluate plasmids?

Prof. Ichizo Kobayashi, University of Tokyo

Dr. Richard Roberts, New England Biolabs

Internal investigators

Wen, Kristie



Other Major Issues and Challenges

- In 2017, most groups used Illumina technology to do *de novo* assembly or align reads to HP26695.

Quite often, they got incomplete assemblies.

- Comparison of the Illumina vs. PacBio results

- How do we check the sequencing quality?

- Resequencing the strains with known sequences

- How do we evaluate these *de novo* assemblies?

- How do we find the unintentional duplicates in the dataset?

- How to evaluate plasmids?

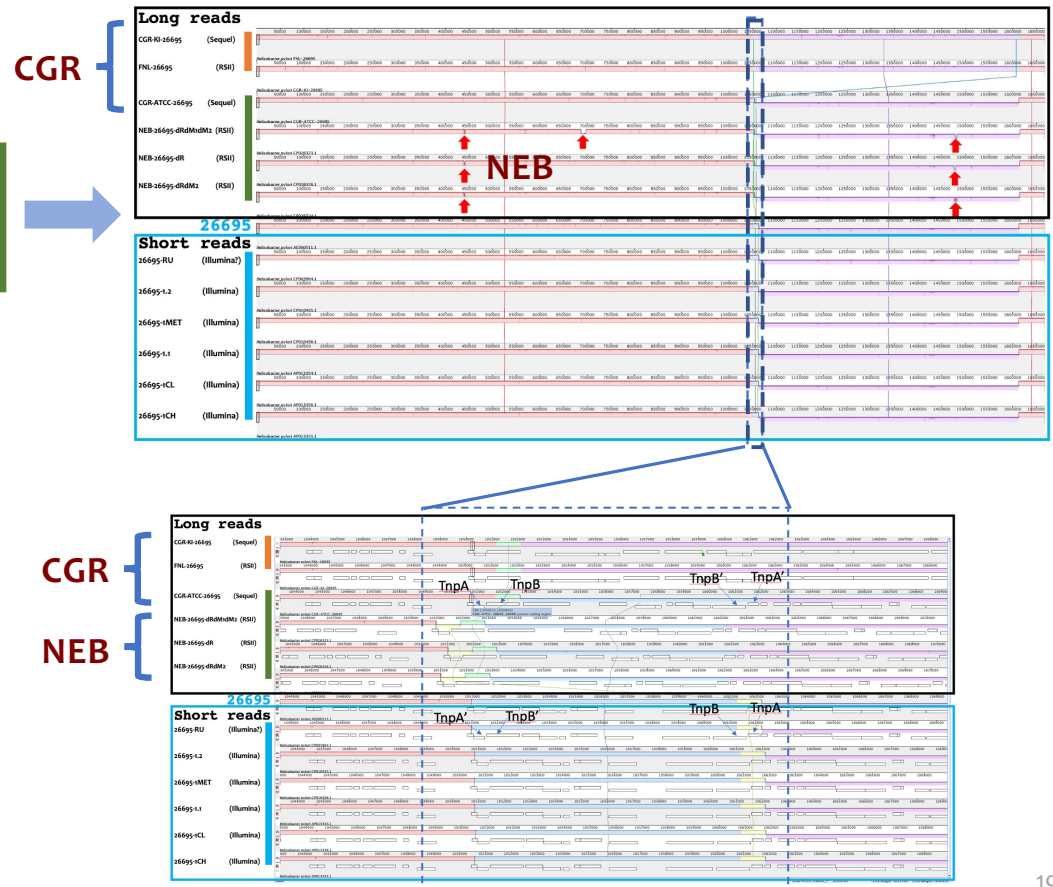
Prof. Ichizo Kobayashi, University of Tokyo

Dr. Richard Roberts, New England Biolabs

Internal investigators

Wen, Kristie

HP26695: Long reads vs. Short reads





Other Major Issues and Challenges

- In 2017, most groups used Illumina technology to do *de novo* assembly or align reads to HP26695.

Quite often, they got incomplete assemblies.

- Comparison of the Illumina vs. PacBio results
- How do we check the sequencing quality?
 - Resequencing the strains with known sequences
- How do we evaluate these *de novo* assemblies?
 - BUSCO score & number of pseudo genes
- How do we find the unintentional duplicates in the dataset?
- How to evaluate plasmids?



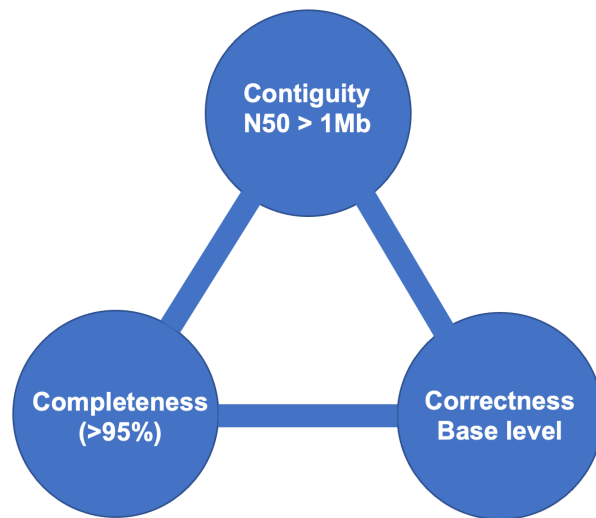
Prof. Ichizo Kobayashi, University of Tokyo

Dr. Richard Roberts, New England Biolabs

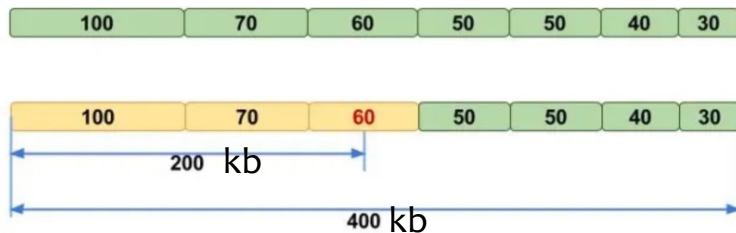
Internal investigators

Wen, Kedest, Kristie

Assessing the Quality of Genome Assemblies with the 3C's



N50 explained



<https://www.pacb.com/blog/beyond-contiguity>

<https://www.molecular ecologist.com/2017/03/29/whats-n50/>

Contiguity is often measured as contig N50, which is the length cutoff for the longest contigs that contain 50% of the total genome length. **In this era of long-read genome assemblies, a contig N50 over 1 Mb is generally considered good.**

Completeness is often measured using BUSCO (Benchmarking Universal Single-Copy Orthologs) scores, which look for the presence or absence of highly conserved genes in an assembly. The aim is to have the highest percentage of genes identified in your assembly, with **a BUSCO complete score above 95% considered good.**

Correctness, the third and final C, is more challenging to measure. **Correctness can be defined as the accuracy of each base pair in the assembly** and is most often measured as concordance of an assembly to a gold standard reference. Of course, when sequencing a novel species there may not be a reference against which to measure. Furthermore, **concordance is only a good measure for accuracy when the gold-standard itself is very high quality and when there is little biological divergence between the reference sample and assembly sample.**

Completeness of Genome Assemblies

Genome analysis

BUSCO: assessing genome assembly and annotation completeness with single-copy orthologs

Felipe A. Simão¹, Robert M. Waterhouse¹, Panagiotis Ioannidis, Evgenia V. Kriventseva and Evgeny M. Zdobnov^{*}

¹Department of Genetic Medicine and Development, University of Geneva Medical School and Swiss Institute of Bioinformatics, rue Michel-Servet 1, 1211 Geneva, Switzerland

We propose intuitive metrics to describe genome, gene set or transcriptome completeness in BUSCO notation - C:complete [D:duplicated], F:fragmented, M:missing, n:number of genes used (Fig. 1). The recovered genes are classified as 'complete' when their lengths are within two standard deviations of the BUSCO group mean length (i.e. within ~95% expectation, Supplementary Fig. S1). 'Complete' genes found with more than one copy are classified as 'duplicated'. These should be rare, as BUSCOs are evolving under single-copy control (Waterhouse *et al.*, 2011), and the recovery of many duplicates may therefore indicate erroneous assembly of haplotypes. Genes only partially recovered are classified as 'fragmented', and genes not recovered are classified as 'missing'. Finally, the 'number of genes used' indicates the resolution and hence is informative of the confidence of these assessments.

Lab resequencing
Or *in silico* polishing?

All passed Contiguity
assessment (N50 > 1Mb)

https://busco-data.ezlab.org/v5/data/lineages/campylobacteriales_odb10.2020-03-06.tar.gz

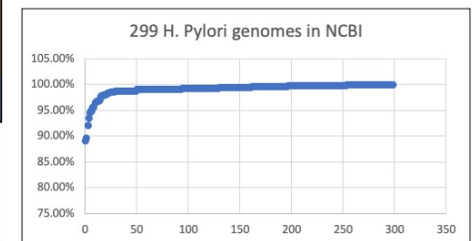
Complete BUSCOs (C)
Complete and duplicated BUSCOs (D)
Complete and single-copy BUSCOs (S)

Fragmented BUSCOs (F)
Missing BUSCOs (M)
Total BUSCO groups searched (n)

AccessionID	C	S	D	F	M	n
CP032043	89.00%	89.00%	0.00%	7.50%	3.50%	628
AP014711	89.60%	89.60%	0.00%	6.70%	3.70%	628
CP032039	92.00%	91.70%	0.30%	5.60%	2.40%	628
CP032027	93.40%	93.20%	0.20%	3.80%	2.80%	628
CP003419	94.40%	94.40%	0.00%	3.70%	1.90%	628
CP023265	94.70%	94.70%	0.00%	3.50%	1.80%	628
CP032046	94.70%	93.90%	0.80%	4.10%	1.20%	628
AP014712	95.40%	95.40%	0.00%	3.50%	1.10%	628
CP006691	95.50%	95.50%	0.00%	3.30%	1.20%	628
CP032041	96.40%	96.20%	0.20%	1.90%	1.70%	628
CP024017	96.50%	96.50%	0.00%	1.80%	1.70%	628
AP014710	96.70%	96.70%	0.00%	2.10%	1.20%	628
CP023266	96.70%	96.70%	0.00%	2.10%	1.20%	628
CP024015	96.70%	96.70%	0.00%	1.90%	1.40%	628
CP032038	96.90%	96.70%	0.20%	2.40%	0.70%	628

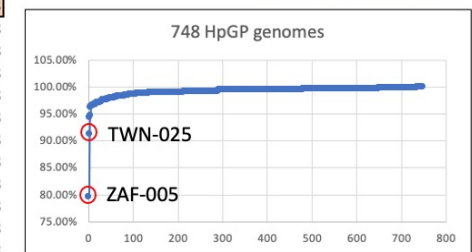
sampleID	C	S	D	F	M	n
ZAF-005	79.50%	79.50%	0.00%	11.00%	9.50%	628
TWN-025	91.20%	90.90%	0.30%	6.20%	2.60%	628
COL-007	94.40%	94.40%	0.00%	2.40%	3.20%	628
COL-304	94.60%	94.60%	0.00%	3.20%	2.20%	628
KOR-035	96.20%	96.20%	0.00%	2.20%	1.60%	628
IND-006	96.30%	96.30%	0.00%	1.40%	2.30%	628
KOR-002	96.30%	96.30%	0.00%	2.70%	1.00%	628
JAP-104	96.50%	96.50%	0.00%	2.10%	1.40%	628
KOR-012	96.50%	96.50%	0.00%	2.70%	0.80%	628
KOR-110	96.50%	96.50%	0.00%	2.40%	1.10%	628
SWE-024	96.50%	96.50%	0.00%	2.10%	1.40%	628
SWE-026	96.50%	96.50%	0.00%	2.40%	1.10%	628
JAP-105	96.70%	96.70%	0.00%	1.80%	1.50%	628
KOR-041	96.70%	96.70%	0.00%	2.10%	1.20%	628
KOR-048	96.70%	96.70%	0.00%	2.10%	1.20%	628

	C	S	D	F	M
min	89.00%	89.00%	0.00%	0.00%	0.10%
median	99.40%	99.40%	0.00%	0.20%	0.30%
max	99.90%	99.80%	3.20%	7.50%	3.70%



GenBank

	C	S	D	F	M
min	79.50%	79.50%	0.00%	0.00%	0.00%
median	99.50%	99.40%	0.00%	0.20%	0.30%
max	100.00%	100.00%	2.10%	11.00%	9.50%



HpGP

Identified 4 genomes need to be improved.

Batch3: 748 strains

299 NCBI GenBank *H. pylori* Genomes (Complete or Chromosome only) (2021-08)

14% genomes with pseudo frameshift genes > 5%

42 out of 299

5%

23

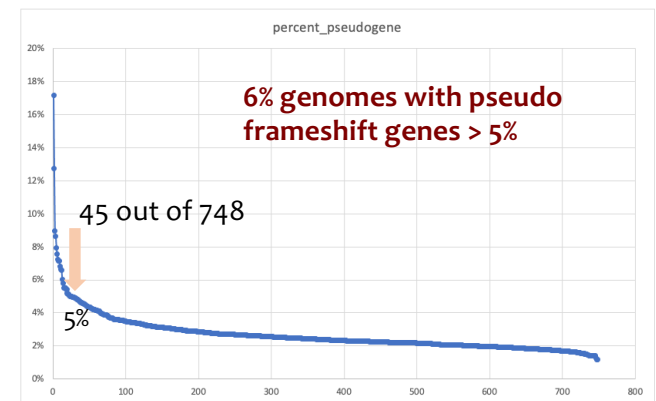
Pseudogenes

748 HpGP Genomes (Batch3)

sampleID	prokka_CDS	prokka_CDS_percent_pseudogene	tRNA	16SrRNA	23SrRNA	length	NCBI_CDS	NCBI_frames	perct
ZAF-005	1926	331	17%	36	2	1650827	NA	NA	
POR-114	1735	221	13%	36	2	1635540	NA	NA	
TWN-025	1670	150	9%	37	2	1625574	1560	249	16%
COG-011	1630	141	9%	36	2	1615291	NA	NA	
SWE-001	1648	131	8%	36	2	1606996	1553	165	11%
KOR-001	1648	125	8%	36	2	1639021	1559	164	11%
SWE-016	1645	119	7%	36	2	1617261	1557	183	12%
COG-005	1608	115	7%	36	2	1615710	NA	NA	
COL-304	1641	117	7%	36	2	1654360	NA	NA	
COL-007	1655	113	7%	36	2	1653789	NA	NA	
PER-016	1661	110	7%	36	2	1689329	NA	NA	
MAL-006	1649	109	7%	36	2	1664187	1596	152	10%
KOR-044	1592	96	6%	36	2	1607983	1541	152	10%
CHI-010	1574	91	6%	36	2	1602017	NA	NA	
KOR-004	1554	86	6%	36	2	1578546	1513	123	8%
SWE-031	1616	89	6%	36	2	1643010	1560	136	9%
KOR-005	1600	88	6%	36	2	1634969	1555	122	8%
KOR-002	1568	85	5%	36	2	1597612	1539	122	8%
KOR-045	1591	86	5%	36	2	1609731	NA	NA	
IND-006	1622	84	5%	36	2	1670468	NA	NA	
KOR-035	1554	80	5%	36	2	1581821	NA	NA	
JAP-002	1543	79	5%	36	2	1562309	NA	NA	
JAP-109	1550	78	5%	36	2	1585338	NA	NA	
KOR-048	1533	77	5%	36	2	1574538	1504	103	7%
MEX-028	1535	77	5%	36	2	1562280	1481	116	8%
KOR-041	1578	79	5%	36	2	1605825	1548	133	9%
SWT-007	1627	81	5%	36	2	1663034	1584	134	8%
KOR-110	1577	78	5%	36	2	1623260	NA	NA	
JAP-104	1578	78	5%	36	2	1629601	NA	NA	
KOR-010	1549	76	5%	36	2	1582866	1513	123	8%
SGP-018	1593	78	5%	36	2	1634713	NA	NA	
SWE-024	1577	77	5%	36	2	1625484	1560	122	8%
KOR-033	1577	76	5%	36	2	1613357	1548	133	9%
COG-004	1560	75	5%	36	2	1615896	NA	NA	
JAP-010	1536	73	5%	36	2	1574397	1516	118	8%
JAP-105	1518	72	5%	36	2	1555420	NA	NA	
MAL-023	1539	72	5%	36	2	1582641	1518	86	6%
SWE-026	1604	75	5%	36	2	1654901	1567	148	9%
SWE-004	1532	71	5%	36	2	1600123	1547	123	
GRE-041	1561	72	5%	36	2	1626901	NA	NA	
SWE-022	1592	73	5%	36	2	1644054	1577	131	
MAL-020	1532	70	5%	36	2	1583392	1513	113	7%
PRI-001	1515	69	5%	36	2	1577958	NA	NA	
SPA-611	1606	73	5%	36	2	1666518	NA	NA	
JAP-111	1533	69	5%	36	2	1563824	NA	NA	
GRE-044	1613	72	4%	36	2	1688552			

45 genomes >= 5%; 17 genomes > 5%

Based on Prokka annotation



tRNA: 36
16SrRNA: 2
23SrRNA: 2
% Pseudo gene: <= 5%

These 17 genomes need to be updated.

Other Major Issues and Challenges

- In 2017, most groups used Illumina technology to do *de novo* assembly or align reads to HP26695.

Quite often, they got incomplete assemblies.

- Comparison of the Illumina vs. PacBio results
- How do we check the sequencing quality?
 - Resequencing the strains with known sequences
- How do we evaluate these *de novo* assemblies?
 - BUSCO score

- How do we find the unintentional duplicates in the dataset?

- Use Ichizo's 20 family strains set

- How to evaluate plasmids?

Prof. Ichizo Kobayashi, University of Tokyo

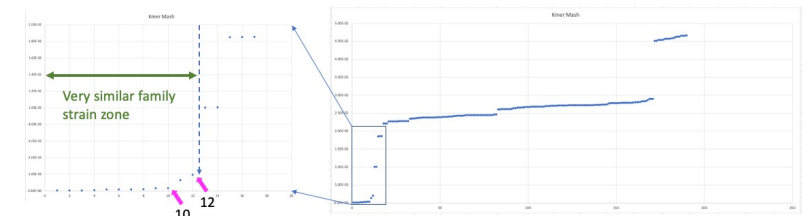
Dr. Richard Roberts, New England Biolabs

Internal investigators

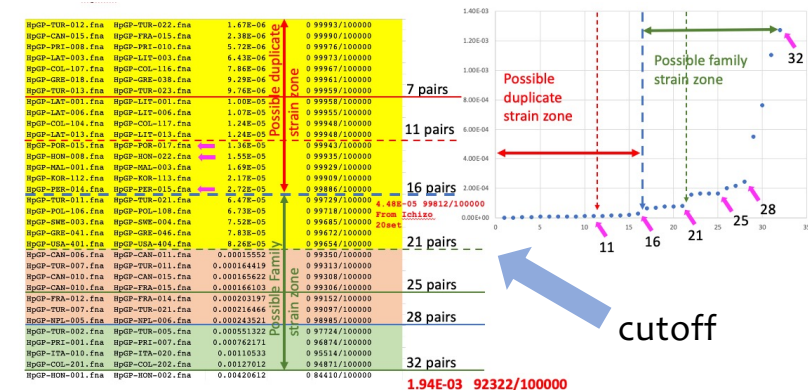
Wen, Kedest, Kristie

Ichizo's 20 family strains mash kmr

K23_chromosomal.fna	K24_chromosomal.fna	4.48E-05	0.99812/100000	K23, K24, K25 three kids from one family	Family strain zone	4.48E-05	99812/100000
K24_chromosomal.fna	K25_chromosomal.fna	4.48E-05	0.99796/100000				
K23_chromosomal.fna	K25_chromosomal.fna	5.80E-05	0.99777/100000				
K133_chromosomal.fna	K134_chromosomal.fna	8.32E-05	0.99642/100000				
K27_chromosomal.fna	K28_chromosomal.fna	0.00011768	0.99341/100000	K27 mem, K28 family kid			
HP94_chromosomal.fna	HP97_chromosomal.fna	0.00016971	0.99294/100000				
HP97_chromosomal.fna	HP98_chromosomal.fna	0.00018199	0.99240/100000	SS 2000/SS 2014			
HP93_chromosomal.fna	HP97_chromosomal.fna	0.00018458	0.99231/100000				
K26_chromosomal.fna	K27_chromosomal.fna	0.00002238	0.98742/100000	K26, K27: from Family K-3 Father / Mother			
K26_chromosomal.fna	K28_chromosomal.fna	0.00011913	0.98673/100000	K26 dad, K28 family kid			
HP79_chromosomal.fna	HP98_chromosomal.fna	0.0012977	0.94764/100000	Kan before / after eradication	10 pairs	1.94E-03	92322/100000
HP77_chromosomal.fna	HP91_chromosomal.fna	0.0019807	0.93252/100000	Kap10000 before / after eradication			
PS7_chromosomal.fna	HP77_chromosomal.fna	0.0100315	0.68074/100000		12 pairs		
PS7_chromosomal.fna	HP92_chromosomal.fna	0.0100953	0.68042/100000				
HP94_chromosomal.fna	K27_chromosomal.fna	0.0184812	0.51225/100000				
HP94_chromosomal.fna	K28_chromosomal.fna	0.0185076	0.51262/100000				
HP94_chromosomal.fna	K26_chromosomal.fna	0.0185384	0.51246/100000				



HpGP strains mash kmr



Other Major Issues and Challenges

- In 2017, most groups used Illumina technology to do de novo assembly or align reads to HP26695.

Quite often, they got incomplete assemblies.

- Comparison of the Illumina vs. PacBio results
- How do we check the sequencing quality?
 - Resequencing the strains with known sequences
- How do we evaluate these de novo assemblies?
 - BUSCO score
- How do we find the unintentional duplicates in the dataset?
 - Use Ichizo's 20 family strains set

- How to evaluate plasmids?

- Exonuclease V digestion screening

Prof. Ichizo Kobayashi, University of Tokyo

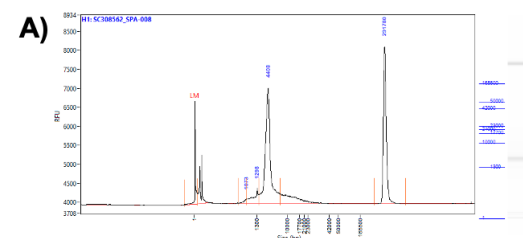
Dr. Richard Roberts, New England Biolabs

Internal investigators

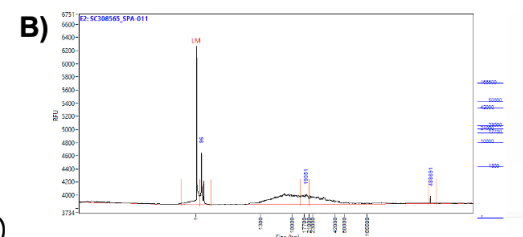
Wen, Kedest, Kristie
Josh Cherry(NCBI), John Dekker(NIAID)

Exonuclease V digestion screening

Status	Sample Number
Exonuclease screening performed	623
Flagged for follow-up assembly	129
Plasmid contig identified in follow-up assembly	52

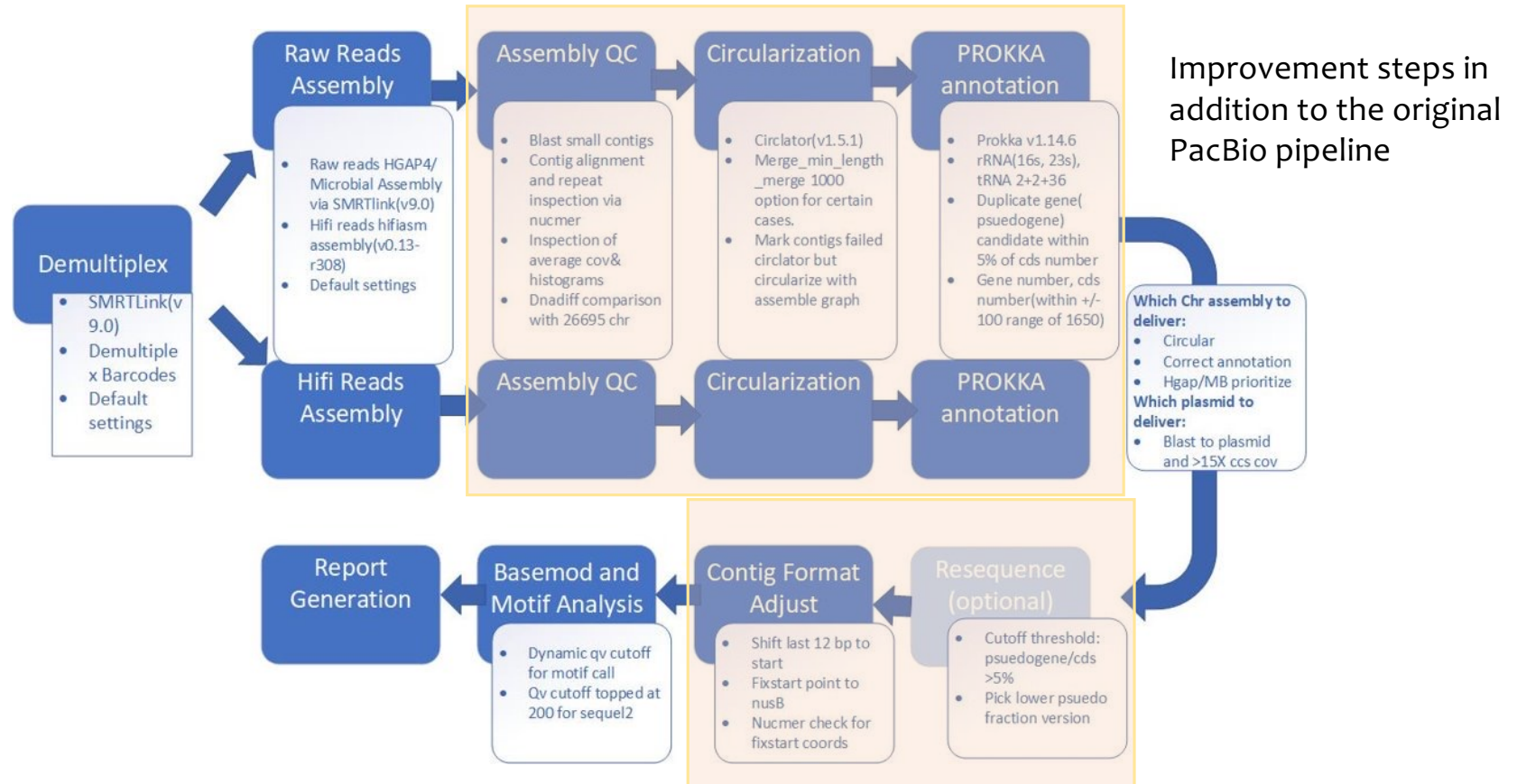


Initially, no plasmid detected.
After screening, plasmids detected



Initially, no plasmid detected.
After screening, no plasmids detected

Improved *H. pylori* Genome Assembly Pipeline



Wen and Kristie



HpGP Data Analyses

Primary Analyses

- Population structure and ancestry (Thorell, Muñoz-Ramirez *et al.*) Landmark paper published.
- Prophages (Vale *et al.*) Under review.
- Methylation profiles (Roberts *et al.*) in preparation.
- Plasmids (Torres R. *et al.*) ongoing.
- Antibiotic resistance to clarithromycin and levofloxacin (Chiner-Oms *et al.*) ongoing.
- Genome- and epigenome wide association analyses of gastric cancer and advanced intestinal metaplasia (Yahara *et al.* → Wang) ongoing.

**HpGP data (1012 genomes)
was finalized in April 2022.**

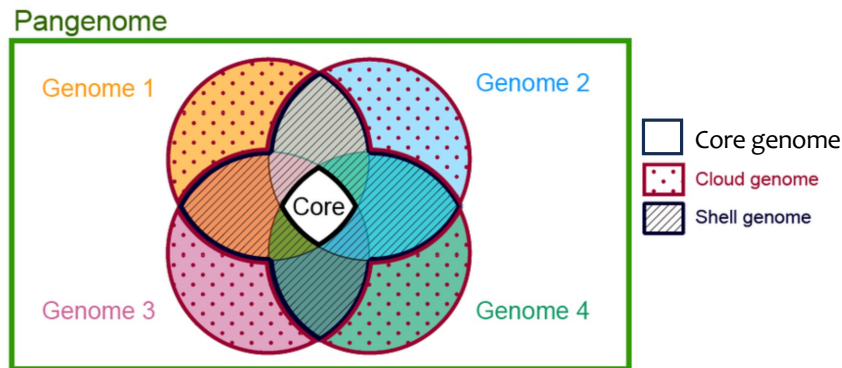
Secondary Analyses

- Rearrangements
- Integrative conjugative elements
- Gene-centric
 - *vacA*, *babA*, *babB*, *babC*, *hopQ*, *cagY*, 16S
- Metabolic pathways
- Mutational signature
- Adaptive differentiation
- Non-coding RNAs

Focus of the *H. pylori* landmark paper

Pan and Core Genome Analyses

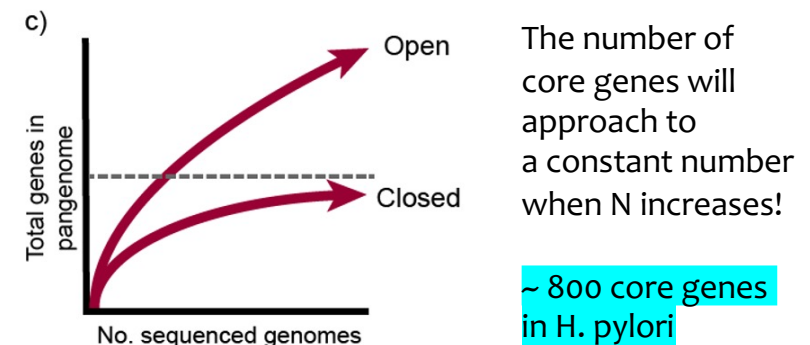
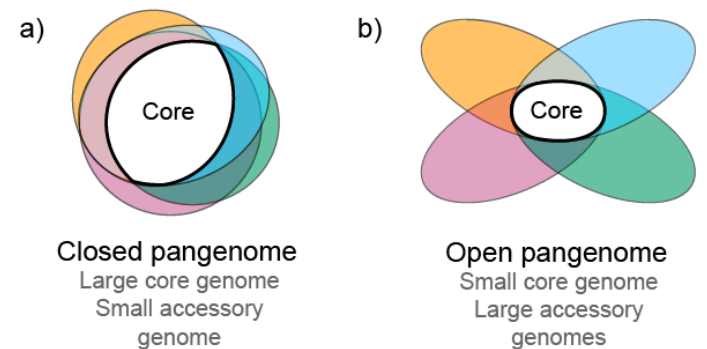
- A pan-genome (pangenome or supragenome) is the entire set of genes from all strains within a clade.



The $\alpha = 0.879 \pm 0.035 < 1$ (> 30 genomes) value of the Heaps' law indicates that **the pan-genome of *H. pylori* is "open"** i.e., the size of the pan-genome tends to diverge when N increases, as concluded in a previous analysis using seven *H. pylori* genomes.

Uchiyama et al. PLoS ONE 11(8): e0159419, (2016)
Tettelin et al. Curr. Opin. Microbiol. 11, 472-477 (2008)
Fischer W et al. NAR, 38, 6089-6101, (2010)

- Open and closed pangenome

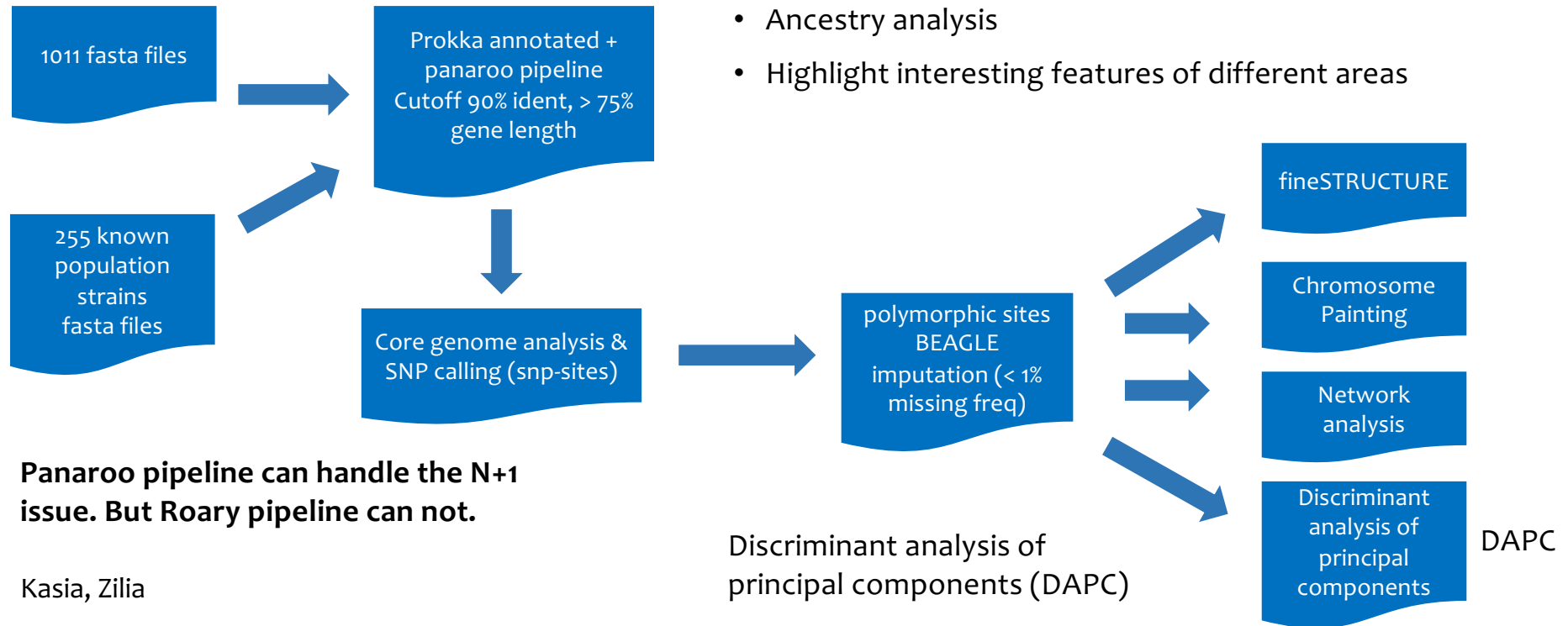


<https://en.wikipedia.org/wiki/Pan-genome>

Focus of the *H. pylori* landmark paper



Analysis workflow



STRUCTURE vs. fineSTRUCTURE

- Data Type:** STRUCTURE works with genotype data, while FineSTRUCTURE operates on phased haplotype data.
- Scope of Analysis:** STRUCTURE is commonly used for inferring population-level admixture, while FineSTRUCTURE is more focused on fine-scale population structure and relationships.
- Output Representation:** STRUCTURE typically presents results as bar plots, while FineSTRUCTURE outputs heatmaps and dendrograms.

In summary, STRUCTURE and FineSTRUCTURE serve similar purposes in terms of population structure analysis, but FineSTRUCTURE provides a more detailed examination of fine-scale genetic relationships by leveraging haplotype information and producing informative visualizations. The choice between the two tools often depends on the specific objectives of the study and the level of resolution required in understanding population structure.

fineSTRUCTURE Results

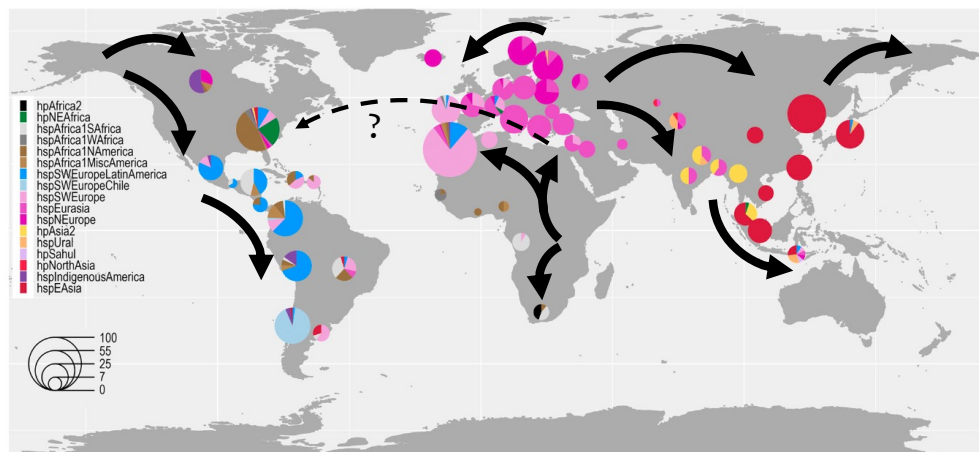


Fig. 1 | World map of HpGP strain origins and population assignments. The area of each pie is proportional to the number of HpGP genomes from each country and colored by the *H. pylori* population (hp) and subpopulation (hsp) as assigned by fineSTRUCTURE (Supplementary Figs. 1 and 2).

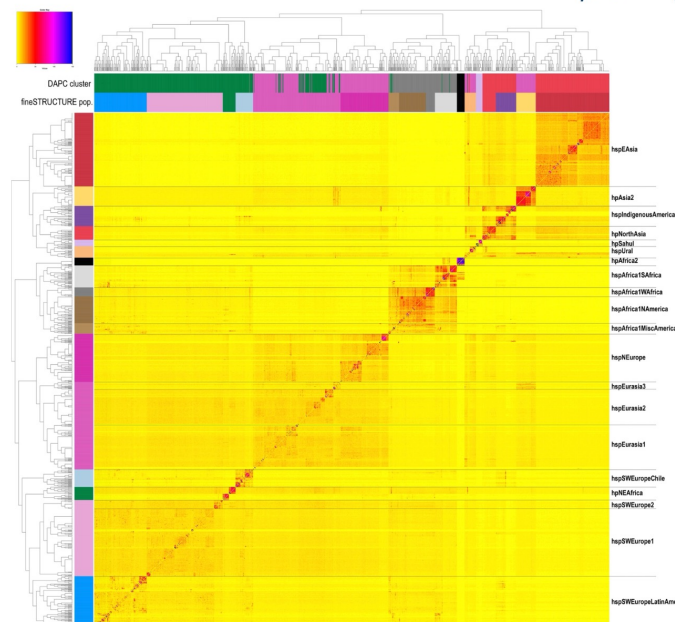
- Revealed four main population clusters
 - (i) Southwest Europe, including Latin America and Northeast Africa
 - (ii) Northern and Central Europe, Middle East, and Central Asia
 - (iii) Western and Southern Africa, including Africa2 and North, South and Central America
 - (iv) North, Central and East Asia, and Indigenous populations in America.

It further formed 17 main subpopulations.

Supplementary Figure 1.

Population structure of global *H. pylori* strains.

The colour of each cell of the matrix indicates the expected number of DNA chunks imported from a donor genome (column) to a recipient genome (row). The inferred tree was generated by Bayesian clustering in fineSTRUCTURE. The colour bars on the top and left indicate suggested *H. pylori* population (hp) and subpopulation (hsp) as in Fig. 1, and the discriminant analysis of principal components, DAPC, K=6 cluster (Fig S3), respectively.



- The hpEurope subpopulations span from the Atlantic coast to South Asia
- Clear differences between East and West Europe. hspNEurope further spans to South Asia.

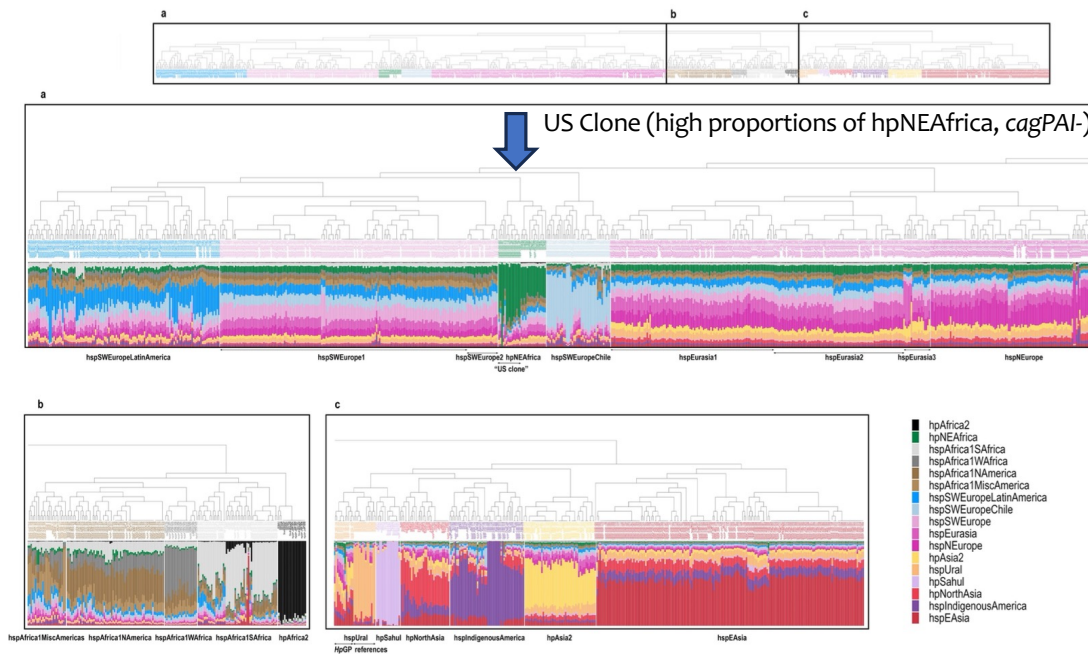
Strains from Eastern Europe and the Middle East are new in this study and have rarely been studied before.

3D plot: https://hpgp.shinyapps.io/Interactive_figures/

Kasia, Zilia

Chromosome Painting Results

Chromosome painting proportions for each genome



Ancestral analysis of Central Asian genomes

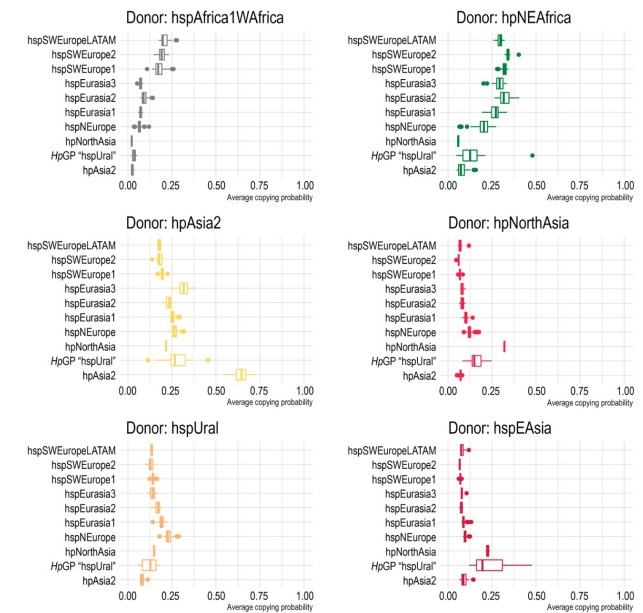
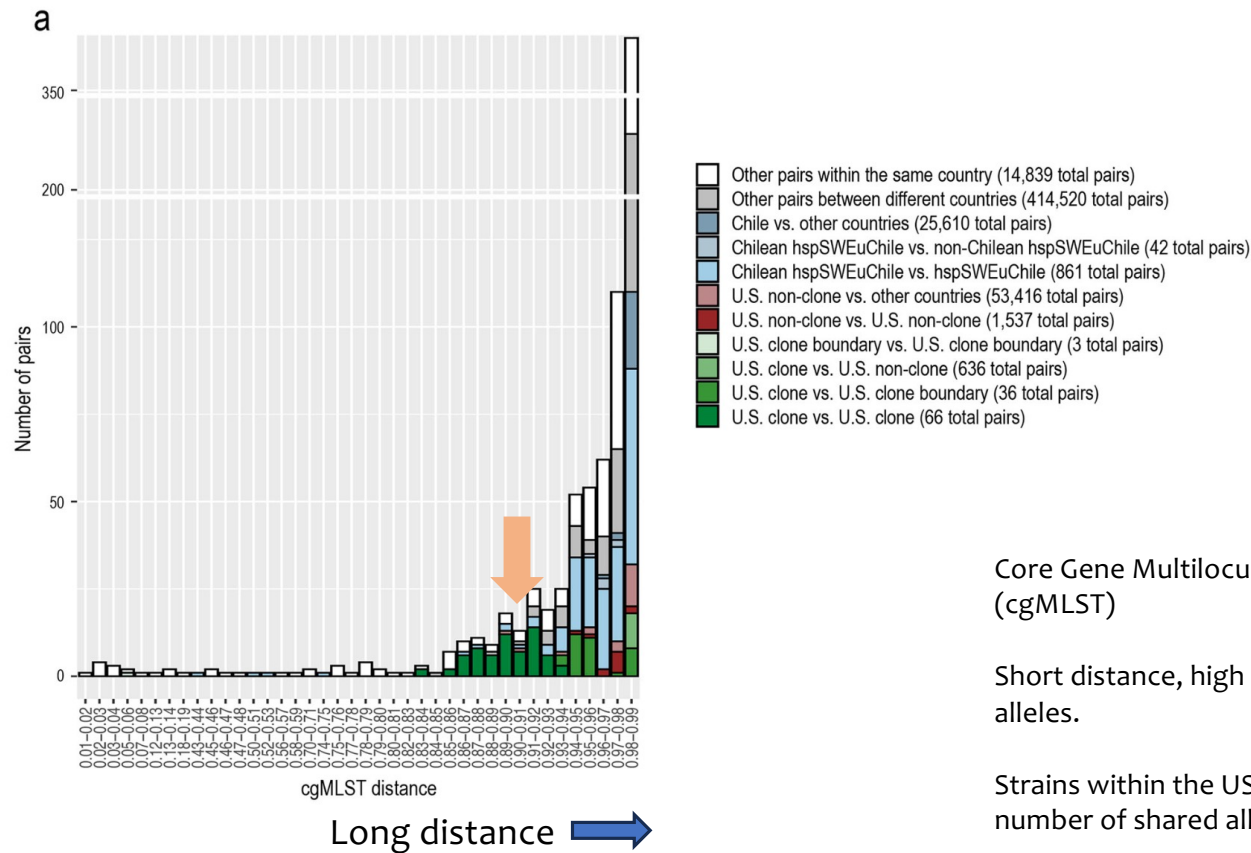


Fig. 3 | Inferred ancestral genomic contributions to the Eurasian HpGP genomes. Ancestral chromosome painting proportions by donor and Eurasian subpopulation. Boxplots show the median value per group, and the 25th and 75th percentiles (hinges), with whiskers extending from the hinge to the largest value no further than 1.5 × IQR (inter-quartile range) from the hinge. Data points beyond the

whiskers are plotted individually. The number of genomes in each respective Eurasian population is hspSWEuropeLatinAmerica, $n = 15$; hspSWEurope2, $n = 12$; hspSWEurope1, $n = 129$; hspEurasia3, $n = 18$; hspEurasia2, $n = 76$; hspEurasia1, $n = 103$; hspNEurope, $n = 95$; hpNorthAsia, $n = 2$; HpGP “hspUral”, $n = 10$; hpAsia2, $n = 27$.

The ancestral contributions to the **central Asian genomes (Eurasian)** confirmed the HpGP “hspUral” clade not to have pronounced contribution by the hspUral references but **relatively high hpAsia2, hpNorthAsia and hspEAsia painting proportions**

In-depth Analysis of the US Deep Clonal Relationships in HpGP



Core Gene Multilocus Sequence Typing (cgMLST)

Short distance, high number of shared alleles.

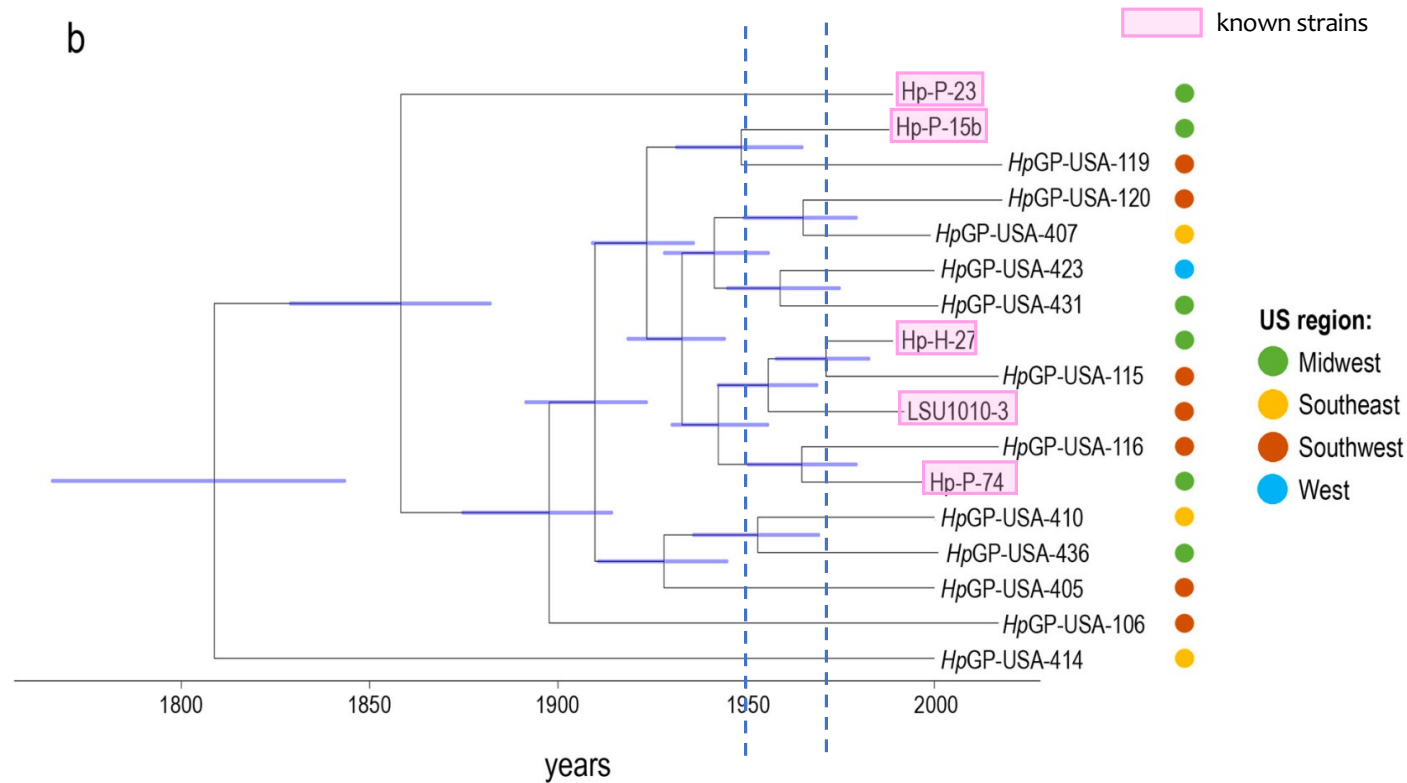
Strains within the US clone have high number of shared alleles.



In-depth Analysis of the US Deep Clonal Relationships in HpGP

ClonalFrameML tree

Estimated it from a common ancestral strain ~ 175 years ago in the US.



Kasia, Zilia

Summary of Population Classifications

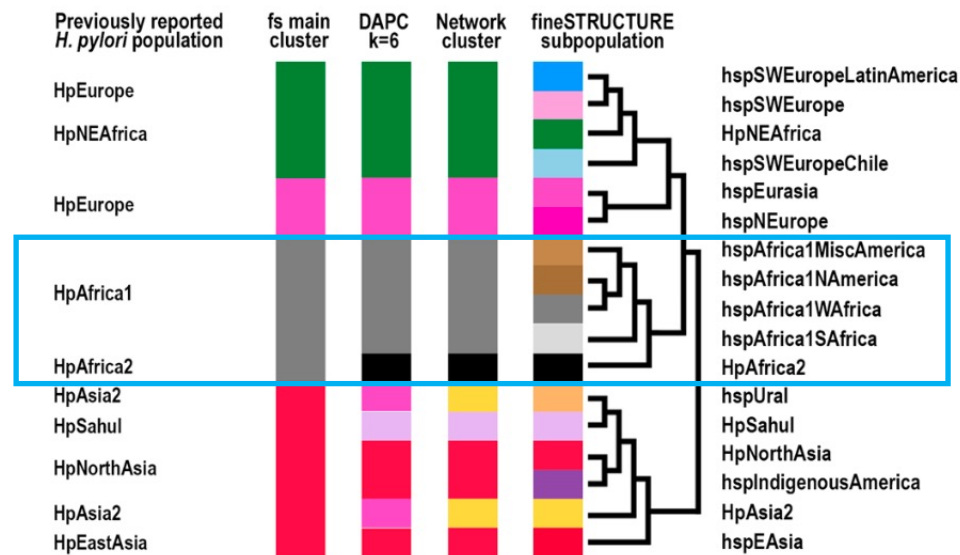
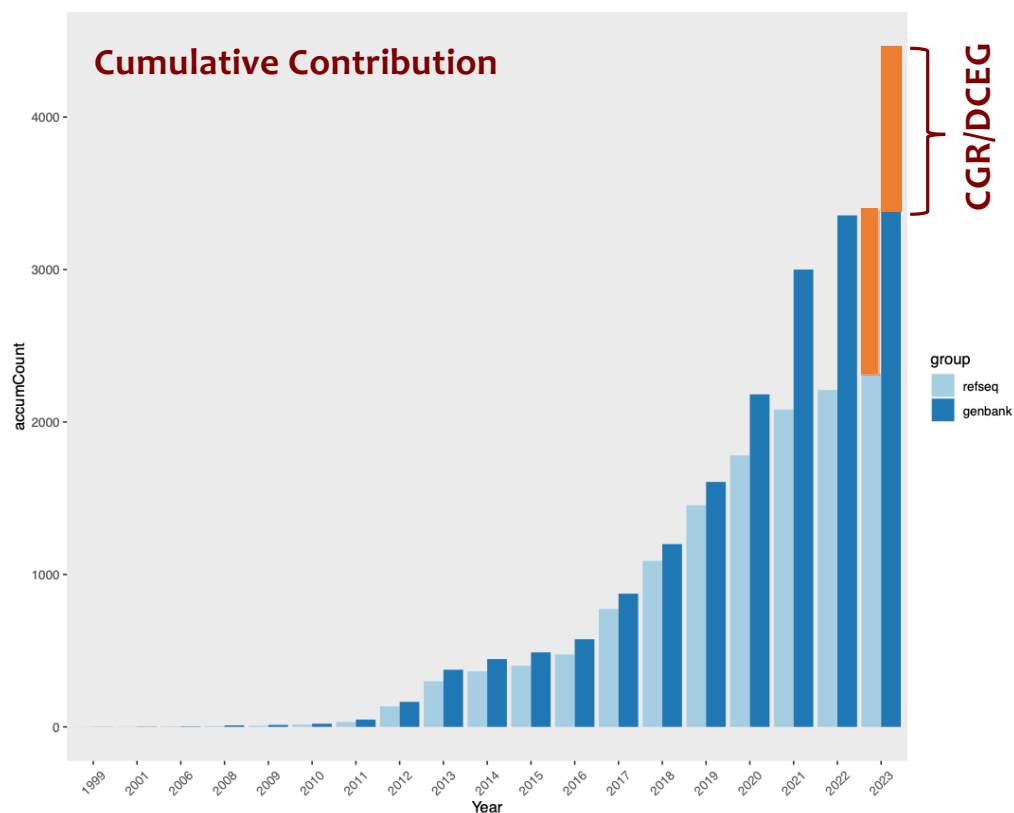


Fig. 5 | Summary of population classifications. Summary of the clustering results using the respective analyses in relation to previously reported MLST and whole genome-based *H. pylori* populations (Hp) and subpopulations (hsp). Colors are based on classifications from the fineSTRUCTURE (fs) analyses visualized in Supplementary Fig. 1, on the $K = 6$ discriminant analysis of principal components, DAPC (Supplementary Fig. 3), and the network clusters (Fig. 2). The topology of the dendrogram to the left is based on the fineSTRUCTURE hierarchical clustering of Supplementary Fig. 1.



Complete *H. Pylori* Genomes in GenBank (2024-01-20)



An official website of the United States government [Here's how you know](#)

NIH National Library of Medicine
National Center for Biotechnology Information

BioProject [Create alert](#) [Advanced](#) [Browse by Project attributes](#)

Display Settings: Send to:

Helicobacter pylori strain:HpGP Accession: PRJNA529500 ID: 529500

Helicobacter pylori Genome Project (HpGP)

The Helicobacter pylori Genome Project (HpGP), an initiative of the U. S. National Cancer Institute (NCI). [More...](#)

[See Genome Information for Helicobacter pylori](#)

[NAVIGATE ACROSS](#)
1076 additional projects are related by organism.

Accession	PRJNA529500
Data Type	Genome sequencing, Epigenomics
Scope	Monoisolate
Organism	Helicobacter pylori [Taxonomy ID: 210] Bacteria; Campylobacterota; Epsilonproteobacteria; Campylobacteriales; Helicobacteraceae; Helicobacter; Helicobacter pylori
Publications	Thorell K <i>et al.</i> , "The Helicobacter pylori Genome Project: insights into H. pylori population structure from analysis of a worldwide collection of complete genomes.", <i>Nat Commun</i> , 2023 Dec 11;14(1):8184
Submission	Registration date: 2-May-2022 Intramural Project NCI HpGP CAS: 10755 - National Cancer Institute
Relevance	Medical

Project Data:

Resource Name	Number of Links
SEQUENCE DATA	
Nucleotide (Genomic DNA)	1012
Protein Sequences	1446355
PUBLICATIONS	
PubMed	1

<https://www.ncbi.nlm.nih.gov/bioproject/PRJNA529500>

<https://zenodo.org/records/10048320>

Acknowledgment



2019

CGR

Sequencing team
Bioinformatics team
Wen, Kristie, Kedest, Yunhu,
Belynda

DCEG

Maria Constanza Camargo
Charles Rabkin
Kai Yu and Bin Zhu



2023

FNLCR

Sequencing and Bioinformatics teams

CGR Research Analysis Support Group

Chad, Sambit, Weiyin and Xin