

Mechanisms of disease

Comparative full-length genome sequence analysis of 14 SARS coronavirus isolates and common mutations associated with putative origins of infection

YiJun Ruan, Chia Lin Wei, Ling Ai Ee, Vinsensius B Vega, Herve Thoreau, Se Thoe Su Yun, Jer-Ming Chia, Patrick Ng, Kuo Ping Chiu, Landri Lim, Zhang Tao, Chan Kwai Peng, Lynette Oon Lin Ean, Ng Mah Lee, Leo Yee Sin, Lisa F P Ng, Ren Ee Chee, Lawrence W Stanton, Philip M Long, Edison T Liu

Summary

Background The cause of severe acute respiratory syndrome (SARS) has been identified as a new coronavirus. Whole genome sequence analysis of various isolates might provide an indication of potential strain differences of this new virus. Moreover, mutation analysis will help to develop effective vaccines.

Methods We sequenced the entire SARS viral genome of cultured isolates from the index case (SIN2500) presenting in Singapore, from three primary contacts (SIN2774, SIN2748, and SIN2677), and one secondary contact (SIN2679). These sequences were compared with the isolates from Canada (TOR2), Hong Kong (CUHK-W1 and HKU39849), Hanoi (URBANI), Guangzhou (GZ01), and Beijing (BJ01, BJ02, BJ03, BJ04).

Findings We identified 129 sequence variations among the 14 isolates, with 16 recurrent variant sequences. Common variant sequences at four loci define two distinct genotypes of the SARS virus. One genotype was linked with infections originating in Hotel M in Hong Kong, the second contained isolates from Hong Kong, Guangzhou, and Beijing with no association with Hotel M ($p < 0.0001$). Moreover, other common sequence variants further distinguished the geographical origins of the isolates, especially between Singapore and Beijing.

Interpretation Despite the recent onset of the SARS epidemic, genetic signatures are emerging that partition the worldwide SARS viral isolates into groups on the basis of contact source history and geography. These signatures can be used to trace sources of infection. In addition, a common variant associated with a non-conservative aminoacid change in the S1 region of the spike protein, suggests that immunological pressures might be starting to influence the evolution of the SARS virus in human populations.

Published online May 9, 2003
<http://image.thelancet.com/extras/03art4454web.pdf>
 See Commentary page

Genome Institute of Singapore, Singapore (Y J Ruan PhD, C L Wei PhD, V B Vega, H Thoreau, J-M Chia, P Ng PhD, K P Chiu PhD, L Lim, T Zhang PhD, L F P Ng PhD, E C Ren PhD, L W Stanton PhD, P M Long PhD, E T Liu MD); **Virology Section, Department of Pathology, Singapore General Hospital, Singapore** (A E Ling MD, S Y Se Thoe PhD, K P Chan MD, L E Oon MD); **Department of Microbiology and Electron Microscopy Unit, National University of Singapore** (M L Ng PhD); **Tan Tock Seng Hospital, Singapore** (S Y Leo MD)

Correspondence to: Dr Edison T Liu, 1 Science Park Road 05-01, Singapore Science Park II, Singapore, 117528 (e-mail: gisliue@nus.edu.sg)

Introduction

The first cases of severe acute respiratory syndrome (SARS) were identified in November, 2002, in Guangdong Province, China. By April, 2003, the epidemic had spread worldwide, affecting 3547 individuals resulting in 182 deaths.¹ In March, 2003, the putative cause of SARS was identified as a new coronavirus.^{2,3} Oropharyngeal specimens from patients with SARS induced a cytopathic effect on Vero E6 tissue culture cells and revealed the presence of coronavirus-like particles. Reverse transcriptase-PCR analysis with random or broadly-specific coronavirus primers amplified a DNA fragment that resembled, but was distinct from, known coronavirus genomes. When tested, these diagnostic PCR methods detected the SARS virus in many clinical samples taken from affected patients. In addition, serological evidence has shown the presence of antibodies specific to the new coronavirus in the serum of patients with SARS.⁴ Collectively, these data strongly implicate the new coronavirus as the cause of SARS, which we designate as SARS-CoV.

SARS-CoV is a member of the Coronaviridae family of enveloped, POSITIVE-STRANDED RNA VIRUSES, which have a broad host range. Some coronavirus infections in people, cattle, and birds cause respiratory disease, whereas other coronavirus infections in rodents, cats, pigs, and cattle lead to enteric disease. The 27–32 kb genomes of coronaviruses, the largest of RNA viruses, encode 23 putative proteins, including four major structural proteins; nucleocapsid (N), spike (S), membrane (M), and small envelope (E). The spike protein, a glycoprotein projection on the viral surface, is crucial for viral attachment and entry into the host cell. In addition, variations of S protein among strains of coronavirus are responsible for host range and tissue tropism.⁵ Differences in virulence of mice coronaviruses have also been linked to genetic variance in the S protein.^{6,7} and the serological response in the host is typically raised against the S protein.⁸ However, the S, M, and N mature proteins all contribute to generating the host immune response as seen in transmissible gastroenteritis coronavirus,⁹ infectious bronchitis virus,^{10,11} pig respiratory coronavirus,¹² and mouse hepatitis virus.¹³

A characteristic of RNA viruses is the high rate of genetic mutation, which leads to evolution of new viral strains and is a mechanism by which viruses escape host defenses. Therefore, from a public-health perspective, understanding the mutation rate of the SARS virus as it spreads through the population is important. Moreover, the genetic mutability of SARS-CoV, especially in the segments encoding the major antigenic proteins, would also have an effect on development of broadly effective vaccines.

GLOSSARY**BLAST**

The basic local alignment search tool is a system for searching similar sequences against all available sequence databases irrespective of whether the query is DNA or protein sequences. The BLAST programs consist of blastn, blastp, tblast, tblastn, tblastx, PSI-BLAST, MEGABLAST, and so on. The most improved point of this software is its high searching speed

CLUSTALW

CLUSTALW is a multiple sequence alignment tool that is commonly used in the bioinformatics community. It produces global multiple sequence alignments through three major phases: pairwise alignment, guide tree construction, and multiple alignment. The guide tree generated by CLUSTALW is an estimate of relations between sequences much like a phylogenetic tree

PAUP*

The Phylogenetic Analysis Using Parsimony (PAUP*) is a well established package for phylogenetic tree construction. It uses various methods, including parsimony, maximum likelihood, and distance methods to estimate phylogenetic relations.

PHRED/PHRAP/CONSED

Bioinformatics tools to assemble shotgun sequences, which are available from University of Washington, Phred=base-calling program; Phrap=assembly program for shotgun sequences; Consed=UNIX-based graphic editor for Phrap sequence assemblies

POSITIVE-STRANDED RNA VIRUSES

Some viruses, such as coronaviruses, carry their genetic material as RNA rather than the more typical DNA-based genomes. Positive-stranded RNA (also called plus-stranded) indicates that the single stranded RNA genome is of the same sense as coding messenger RNA

We aimed to determine the complete genome sequence for five SARS-CoV related isolates from a single SARS index case, three associated primary contact cases, and one secondary contact and compare them with nine other SARS-CoV isolates available in public-domain databases.

Methods**Patients**

We obtained five positive isolates for coronavirus from the index patient of the outbreak (SIN2500), three primary contacts of this patient (SIN2774, SIN2748, SIN2677), and one secondary contact (SIN2679) who was related to the index patient, but contracted SARS from another primary contact not included in this study. All patients fitted the WHO case definition for probable SARS¹⁴—fever of 38°C or higher, respiratory symptoms (eg, cough, shortness of breath, difficulty in breathing), hypoxia and chest radiograph changes suggestive of pneumonia, and history of close contact with another patient with SARS or travel to a region with documented community transmission of SARS within 10 days of onset of symptoms. The index patient had a history of travel to Hong Kong and had stayed in hotel M.¹⁵ The viral sources were all from respiratory samples: two endotracheal tube aspirates, two nasopharyngeal aspirates, and one throat swab obtained from the patients between 0 and 11 days after onset of symptoms.

Procedures

The virus-containing samples were inoculated into various cell lines including Vero cells, which showed a cytopathic effect characterised by generalised rounding against a granular background seen 5–11 days after inoculation. We maintained the cells at 33°C and repassaged them after 7 days of incubation.

We spotted washed cell pellets from harvested cells showing cytopathic effects onto glass slides and overlaid

them with acute and convalescent serum samples from the patients from whom we obtained the respiratory samples. All five cell lines showed reactivity by immunofluorescence (methods available from authors) with convalescent serum samples from the respective patients. We confirmed the presence of the SARS coronavirus by PCR on the cell supernatants using SARS-specific primers.^{2,4} When tested, two of the infected Vero cell lines also had coronavirus-like particles on electron microscopy. We isolated the viral RNA template from the supernatants of Vero cells that showed cytopathic effects by centrifugation at 23 000 relative centrifugal force for 2.5 h to pellet the viral particles and extracted RNA with the QiAmp viral RNA mini kit (Qiagen, Valencia, CA, USA).

To sequence the viral genome we used both shot-gun and specific priming approaches. From the RNA templates, we synthesised double strand cDNA with the SuperScript cDNA system (Invitrogen, Carlsbad, CA, USA). The cDNA were PCR amplified (20 cycles) with random 16-mer by Platinum Taq polymerase. The amplicons were then cloned into pCR2.1-TOPO vector (Invitrogen). We selected random clones for single pass sequencing analysis on ABI3730 sequencers (Applied Biosystems, Foster City, CA, USA). We compared sequence data generated from the library with human, mouse, and viral genome databases managed at the US National Center for Biotechnology Information (NCBI) site (<http://www.ncbi.nlm.nih.gov>). Since previous sequences from reverse transcriptase-PCR of SARS patient samples are most closely matched to the mouse hepatitis virus at protein level, we used the mouse

Genome sequences	Accession number
SARS-CoV isolates	
<i>Complete</i>	
SIN2500	AY283794
SIN2677	AY283795
SIN2679	AY283796
SIN2748	AY283797
SIN2774	AY283798
TOR2	NC_004718
URBANI	AY278741
CUHKU-W1	AY278554
HKU-39849	AY278491
<i>Partial</i>	
GZ01	AY278489
BJ01	AY278488
BJ02	AY278487
BJ03	AY278490
BJ04	AY279354
Coronavirus isolates	
MHV strain 2	AF201929.1
MHV Penn 97-1	AF208066
MHV	NC_001846
MHV strain ML-10	AF208067
Bovine CoV	NC_003045.1
BCoV Quebec strain	AF220295.1
AIBV	NC_001451.1
TGV	NC_002306.2
TGV	Z34093.1
HCoV 229E	NC_002645.1
PEDV strain C	AF353511.1
PEDV	NC_003436.1
Rat SDAV	AF207551.1
Porcine HEV	AY078417.1

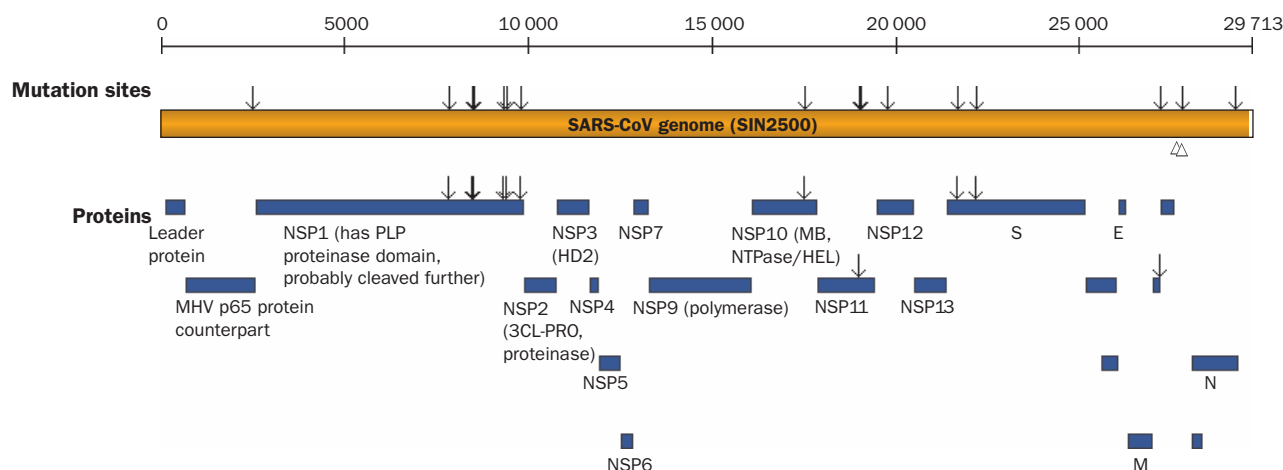


Figure 1: **Genome structure of SARS-CoV**

NSP=Non-structural proteins. S=spike protein. E=small envelop protein. N=nucleocapsid. M=Membrane protein. MHV=murine hepatitis virus. MD=metal binding. 3CL-PRO=3C-like proteinase. The top scale shows the approximate nucleotide position along SARS-CoV genome (golden bar) determined from the Singapore isolate SIN2500. The arrows on top of the genome bar map the locations where nucleotide sequence variations, present in two or more isolates, were detected. The two open triangles point to the locations where the two multi-nucleotide deletions occurred. The SARS genome is predicted to encode 23 putative mature proteins (blue bars). The arrows on top of the protein bars indicate the location of aminoacid changes.

hepatitis virus genome sequence (NC_001846) as a backbone to align the viral cDNA sequence fragments for positions on the viral genome. For sequence positioning we also used limited unpublished SARS viral sequences (EMC1 to EMC8) posted April 5, 2003 on a secure website at WHO.¹⁶ We designed sequence-specific primers on the basis of cDNA sequence data to fill the gaps of 1–2 kb distance. After we completed the first rough genome sequence for the sample SIN2500, we designed 30 primer pairs incorporating sequence information from the newly available TOR2 SARS virus¹⁷ to cover the whole viral genome with each primer pair amplifying about 1200 base pairs. We then used this sequence-specific priming approach to generate reverse transcriptase-PCR fragments for the five whole viral genomes. Each PCR fragment was directly sequenced from both directions inward and outward, in duplicate. Therefore, for any region of the genome there was six to eight-fold coverage. We used the PHRED/PHRAP/CONSED package (University of Washington, Seattle, WA, USA; <http://www.phred.org>) to process all the raw sequence reads for base calling, assembly, and editing. Nucleotide differences in the assembled genome sequences were also checked manually for accuracy. Sequence regions that were poor quality were resequenced either from

purified PCR fragments or cloned plasmid. All genetic variations of Singapore isolates identified when compared with available SARS-CoV genome sequences were further confirmed by primer extension genotyping technology (Sequenom, San Diego, CA, USA). The panel shows the SARS-CoV isolates we accessed from GenBank. We determined the locations of point mutations by aligning the 14 SARS sequences with CLUSTALW.¹⁸ We calculated putative coding sequences by completing the multi-alignment of the 14 SARS sequences with the TOR2 nucleotide sequence annotation as the reference. The annotations of the proteins are taken from the corresponding entries in the NCBI Entrez site.

Associations between the members of the coronaviridae family to the SARS virus were assessed by comparing overlapping fragments of the SIN2500 genomic sequence against a database of coronavirus sequences. To calculate regional homologies with sister coronaviruses we used sequences within sliding 200 base pair windows sampled in a tiled fashion every 100 base pairs across the viral genomes using the BLAST BLASTN program from WU-BLAST (Washington University, St Louis, MO, USA).¹⁹ In the heat map that was generated, the SARS genomic fragments are plotted along the horizontal axis in the order they appear in

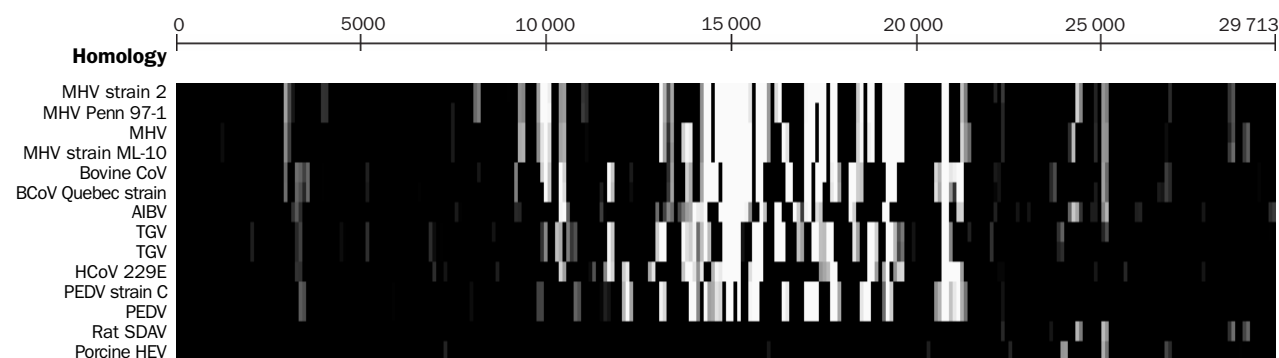


Figure 2: **Homology of SARS-CoV genome sequence (SIN2500) to other coronaviruses**

A heat map created by comparing overlapping fragments of the SARS-CoV genome sequence against a database of coronavirus sequences. The SARS fragments are plotted along the horizontal axis in the order they appear in the genome, and the other coronaviruses are plotted vertically. The brightness of a pixel corresponds to the strength of the match between a SARS fragment and a coronavirus genome; the smaller the p value, the lighter the pixel.

the genome, and the corresponding fragments from other coronaviruses are placed vertically. The brightness of a pixel corresponds to the strength of the match between a SARS fragment and a coronavirus genome; the smaller the p value, the brighter the pixel. At $p=1$ it is black, and the brightness is proportional to $\log(1/p)$ until p is less than 10^{-11} , when it is white. The panel shows the accession numbers of the coronavirus sequences used.

Statistical analysis

In the analysis of the common sequence variations, the probability of co-occurrence of multiple polymorphisms or mutations in an isolate was used as a measure of significance. We used 13 of the samples (we excluded sample BJ04 because of substantial missing sequence information) and restricted our attention to 26 140 loci at which nucleotides were determined in all 13 samples. The null hypothesis was that the nucleotides at these loci were obtained by mutating a single consensus sequence independently at random at each position of each sequence. The mutation rate was estimated from the data by the fraction of positions in the various genomes that differed from the consensus sequence obtained by taking the most frequent nucleotide at each position. We also tested a weaker null hypothesis, in which the mutations taking place at any given locus in different samples are independent, but arbitrary dependence between the loci is possible subject to this constraint. Details of the analytical approach are described in webappendix 1 (<http://image.thelancet.com/extras/03art4454webappendix1.pdf>) and on our website (<http://www.gis.astar.edu.sg/homepage/toolssup.jsp>).

Phylogenetic analysis of the SARS viral genomes was done with PAUP*²⁰ with the maximum probability criterion. We used the default variable settings, with one exception: the substitution rates were estimated from the data. A separate phylogenetic analysis done with CLUSTALW¹⁸ gave the same structure.

Role of the funding source

The sponsor of the study had no role in study design, data collection, data analysis, data interpretation, or in the writing of the report.

Results

We have sequenced the complete genomes of SARS-CoV from five Singapore isolates derived from one index case (SIN2500), three primary cases (SIN2677, SIN2748, and SIN2774), and one secondary case (SIN2679). These sequences showed that the genomes of SARS-CoV isolated in Singapore are comprised of 29 711 bases, with the exception of a five-nucleotide deletion in strain SIN2748 and a six-nucleotide deletion in SIN2677. Initial BLAST analysis suggested that the Singapore SARS virus is similar to, but distinct from, the group 2 coronaviruses in the Coronaviridae family of enveloped and positive-stranded RNA viruses. As in the recently sequenced SARS-CoV (HKU39849, CUHK-W1, TOR2, and URBANI), the Singapore SARS virus contains 11 predicted open reading frames that encode 23 putative mature proteins with known and unknown functions. Most of the non-structural proteins seem to be encoded in the first half of the genome including nspl and nsp2, with putative proteinase function and nsp9 RNA-dependent RNA polymerase, whereas most of the structural proteins such as spike, membrane, envelop, and nucleocapsid are located in the second half of the genome (figure 1, webfigure: <http://image.thelancet.com/extras/03art4454webfigure.pdf>). The haemagglutinin esterase, which is common in the group 2 coronavirus, is missing in the SARS-CoV genome, suggesting that some of the non-structural genes are dispensable in coronaviruses.

Although the genome organisation of SARS-CoV is similar to that of other coronaviruses, SARS-CoV is only distantly related to any coronavirus member, irrespective of species specificity, at both nucleotide and aminoacid levels (figure 2). In assessing the homology with coronaviridae genomes from other species with the SARS-CoV, sequence

Singapore cases										Overseas cases						Variations frequency	ORF/protein	AA change (SIN2500→variation)
Index case SIN2500 29 711	Primary SIN2774 29 711	Secondary SIN2679 29 711	Hanoi URBANI 29 727	Hong Kong CUHKW1 29 712	N China BJ01 28 920	N China BJ03 29 291	N China BJ04 24 774	N China BJ02 29 430	N China BJ04 24 774	Genome position	Primary SIN2748 29 706	Primary SIN2677 29 705	Canada TOR2 29 712	Hong Kong HKU39849 29 727	S China GZ01 29 429	N China BJ02 29 430	N China BJ04 24 774	
2601	T	T	T	T	T	T	T	T	T	2	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Silent
7919	C	C	C	C	C	C	C	C	C	2	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	A→V
8559	T	T	T	T	T	T	T	T	T	2	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Silent
8572	G	G	G	G	G	G	G	G	G	2	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	V→L
9404	T	T	T	T	T	T	T	T	T	5	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	V→A
9479	T	T	T	T	T	T	T	T	T	2	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	V→A
9854	C	C	C	C	C	C	C	C	C	3	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	Orf1a (nsp1)	A→V
17564	T	T	T	T	T	T	T	T	T	6	Orf1ab (nsp10, helicase)	Orf1ab (nsp10, helicase)	Orf1ab (nsp10, helicase)	Orf1ab (nsp10, helicase)	Orf1ab (nsp10, helicase)	Orf1ab (nsp10, helicase)	Orf1ab (nsp10, helicase)	D→E
19064	A	A	A	A	A	A	A	A	A	2	Orf1ab (nsp11)	Orf1ab (nsp11)	Orf1ab (nsp11)	Orf1ab (nsp11)	Orf1ab (nsp11)	Orf1ab (nsp11)	Orf1ab (nsp11)	Silent
19084	T	T	T	T	T	T	T	T	T	4	Orf1ab (nsp11)	Orf1ab (nsp11)	Orf1ab (nsp11)	Orf1ab (nsp11)	Orf1ab (nsp11)	Orf1ab (nsp11)	Orf1ab (nsp11)	I→T
19838	A	A	A	A	A	A	A	A	A	4	Orf1ab (nsp12)	Orf1ab (nsp12)	Orf1ab (nsp12)	Orf1ab (nsp12)	Orf1ab (nsp12)	Orf1ab (nsp12)	Orf1ab (nsp12)	Silent
21721	G	G	G	G	G	G	G	G	G	2	Spike glycoprotein	Spike glycoprotein	Spike glycoprotein	Spike glycoprotein	Spike glycoprotein	Spike glycoprotein	Spike glycoprotein	G→D
22222	T	T	T	T	T	T	T	T	T	5	Spike glycoprotein	Spike glycoprotein	Spike glycoprotein	Spike glycoprotein	Spike glycoprotein	Spike glycoprotein	Spike glycoprotein	I→T
27243	C	C	C	C	C	C	C	C	C	3	Putative protein	Putative protein	Putative protein	Putative protein	Putative protein	Putative protein	Putative protein	P→L
27827	T	T	T	T	T	T	T	T	T	6	Non-coding	Non-coding	Non-coding	Non-coding	Non-coding	Non-coding	Non-coding	
29279	A	A	A	A	A	A	A	A	A	2	Nucleocapsid	Nucleocapsid	Nucleocapsid	Nucleocapsid	Nucleocapsid	Nucleocapsid	Nucleocapsid	Q→P
Linked to Hotel M										No link to Hotel M								

Figure 3: Sequence comparisons of the 14 available SARS-CoV genomes

Only those variant sequences (red) that were present in at least two independent sequences are shown. See webappendix2 for complete list of all variant nucleotides. The frequency of appearance of each variant nucleotide is presented. Highlighted in yellow are four sequence positions that define two distinct genotypic variants of SARS-CoV. The position of each nucleotide is based upon the URBANI SARS-CoV sequence and the corresponding encoded protein or uncharacterised open reading frames (ORF) are indicated. The effect, if any, on the encoded aminoacid (AA) is described. Nucleotides that are missing (-) or ambiguous (N) in the genome sequences are indicated and the length of each genome sequence is given. The patients' countries of origin and association of patient with Hotel M are provided for all viral genome sequences, and the relative order of transmission is provided for the Singapore cases.

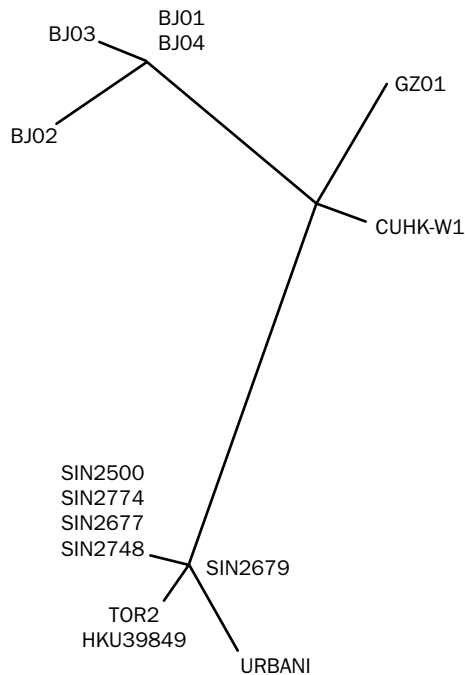


Figure 4: **Molecular relations between the 14 SARS-CoV isolates**

Phylogenetic trees obtained by applying PAUP* to complete genome sequences of all 14 SARS-CoV samples. The tree was built with only sequence variants that occurred at least twice.

conservation seems to be restricted to the middle part of the genome, between bases 14 000 and 21 000, where the RNA-dependent polymerase and several uncharacterised proteins (eg, orf1ab:nsp 10–13) are located. The remainder of the genome, especially at the 5' and 3' regions, diverges from other strains at the nucleotide and aminoacid level.

We aligned the five complete genome sequences of Singapore SARS-CoV isolates and the nine SARS-CoV genomes from outside of Singapore, which were sequenced by others, and investigated the genetic variations between these 14 genomes (webappendix 1). In total, there were 127 single nucleotide sequence variations, one deletion of six nucleotides (nt 27782–27787) in strain SIN2677, and one deletion of five nucleotides (nt 27810–27814) in strain SIN2748 (webappendix 2, <http://image.thelancet.com/extras/03art4454webappendix2.pdf>). Both these deletion sites were in the non-coding sequences between an uncharacterised open reading frame and the nucleocapsid protein. Of the 127 base substitutions, 94 changed the aminoacid sequence (webappendix 2). The mutations were in the following open reading frames: orf1a polyprotein, orf1a RNA-polymerase, orf1ab: nsp10 to nsp 13, spike glycoprotein, membrane nucleocapsid, and several uncharacterised putative proteins.

To eliminate mutational noise induced in culture from real strain differences, we reanalysed the data using a probabilistic approach. Mutations that might have been artifacts of cell culture would occur only once in our survey, whereas sequence variants associated with common ancestry should be seen in multiple isolates. Of the 127 sequence variations in the 14 isolates, 16 variant loci were identified in two or more isolates, and eight were seen in three or more isolates (figure 3). With the more stringent criterion, four loci recurred five or more times in the 14 SARS-CoVs analysed: C/T polymorphisms at position 9404 resulting in a valine to alanine change in orf1a (nsp1); position 19 084 leading to an isoleucine to threonine change in orf1ab (nsp11); position 22 222 changing an isoleucine residue to threonine in the S1 portion of the spike protein, and position 27 827 which is in a non-coding region (figure 3). In addition, a T/G polymorphism is noted at position 17 564 changing orf1ab (nsp10, helicase domain) from an aspartic acid to a glutamic acid. Sequence variants at these four loci segregate together as a specific genotype. For example, isolates CUHKW1, GZD1, BJ01, BJ02, BJ03, BJ04 all have the configuration C:G:C:C at those nucleotide positions, whereas isolates SIN2500, SIN2774, SIN2679, SIN2677, TOR2, URBANI, HKU39849 have the configuration T:T:T:T (figure 3). Assuming that all base substitutions were random events propagated in the Vero cells, the probability of four specific nucleotide changes occurring concurrently is very low. The significance is at $p < 10^{-60}$ when the null hypothesis is that each locus in each sample mutates independently, and $p < 10^{-15}$ when dependence is allowed among the loci. The C:G:C:C and T:T:T:T genotypes are, therefore, very unlikely to have emerged by chance, and might be evidence for the first genetic signature of strain differences in the SARS virus. All isolates with the

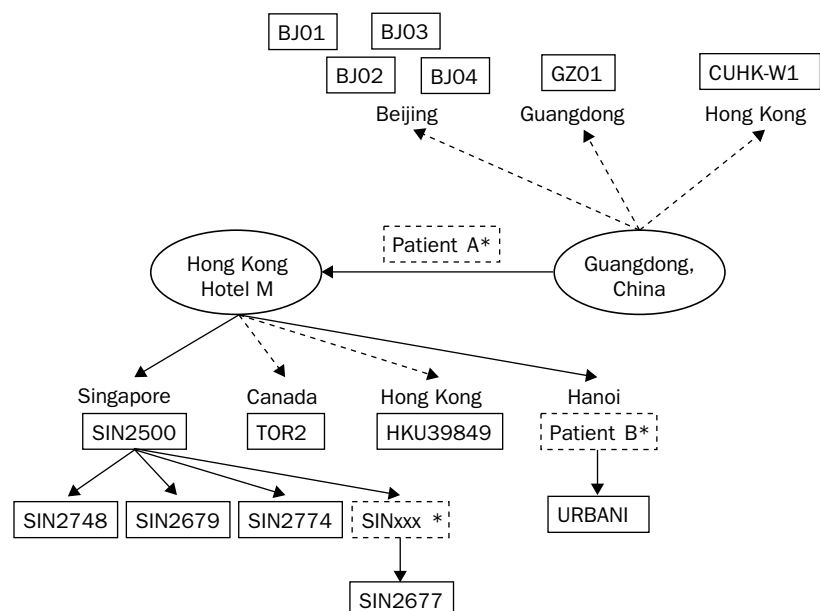


Figure 5: **Clinical relations between the 14 SARS-CoV isolates**

The routes of transmissions are shown for the 14 viral isolates that have been sequenced, indicated with solid boxes. Solid arrows=routes of transmission from Hotel M are known to be direct. Broken arrow=direct relation information is not available. Patient A is the Hong Kong index patient who travelled from Guangdong to Hong Kong and transmitted SARS to others at Hotel M, who then travelled to Singapore, Canada, and Vietnam, thus becoming index cases in those countries.¹⁵ The routes of transmission from Guangdong are unknown and shown as dotted arrows. *Dashed boxes are routes of transmission that are uncertain.

T:T:T:T genotype were linked to infection acquired at the Hotel M in Hong Kong,¹⁸ whereas none of the C:G:C:C genotype isolates had this association (figure 3). Phylogenetic analysis based on the common variant sequences (defined as present in two or more isolates) confirmed that the cases associated with exposure in the Hotel M formed a cluster that was distinct from the other isolates (figure 4). On the basis of the molecular and contact history, we have reconstructed a probable lineage map of the SARS-CoV infections investigated here (figure 5).

In addition to this four-locus genotype, four other common variant sequences (all occurring three to four times in the 14 isolates) seem to further define subgroups geographically (figure 3). The variant sequences at position 19 084 seem to distinguish the Singaporean isolates from all others. Although all nine isolates from outside Singapore showed a C at this position, four of the five isolates from Singapore had a T. The only reversion from T to a C in SIN2679 (a secondary contact case) might be the result of a back-mutational event potentially occurring during the passage of the virus. Taking into account missing data, the polymorphisms at nucleotide positions 9854, 19 838, and 27 243 all segregate with isolates identified specifically in Beijing. Thus, these common sequence polymorphisms might be useful in identifying the differential source of a SARS viral infection.

Discussion

Although the genome organisation of SARS-CoV is similar to that of other coronaviruses, the SARS-CoV sequence is only distantly related to any coronavirus member. Results of earlier reports⁴ suggested that SARS-CoV more closely resembled the cow coronavirus and the mouse hepatitis virus by comparing a conserved 215 aminoacid segment of the polymerase protein product. However, when taking into account the entire genomic sequence, the strength of the associations was reduced, confirming the reports of others that the SARS coronavirus is a completely new pathogenic strain that does not arise from a simple recombination of known existing strains.^{2,21}

Since the S1 subunit of the spike protein is the major antigenic moiety for coronaviruses and is not an essential structural protein, it is prone to high mutation rates as the virus evolves in host populations. That the S1 region did not seem to have excessive numbers of base substitutions suggests that the viral isolates have not been subject to immunological selection.^{22,23} However, because all samples were from viral cultures propagated in Vero cells, some, if not most, of these 129 mutations might have occurred during in-vitro expansion and not because of host pressures.²⁴ In addition, some of the available nucleotide substitutions might have been the result of sequencing errors since several of the sequences were submitted in draft form. To reduce the effects of these technical artifacts, we restricted our analysis to the 16 loci with recurrent mutations among the 14 isolates. These loci are the sequence variants most likely to have been resident in human populations.

Of special interest are the nucleotide changes in four of these loci (positions 9404, 19 084, 22 222, and 27 827) that recurred five or more times. The base substitutions at these locations are highly restricted and segregate together as specific genotypes (C:G:C:C *vs* T:T:T:T). Thus, it is highly unlikely that the C:G:C:C and T:T:T:T genotypes emerged by chance. Rather, we believe this to be evidence for the first genetic signature

of strain differences in the SARS virus. All isolates with the T:T:T:T genotype were linked to infection acquired at the Hotel M in Hong Kong,¹⁵ whereas none of the C:G:C:C genotype isolates had this association.

The index case from Singapore (from which the Singaporean infections described herein were derived) acquired the SARS-CoV infection while staying at Hotel M. The TOR2 virus, cultured in Canada, and the HKU39849 isolate from Hong Kong, were from patients who became infected through contact at Hotel M, although perhaps not directly. The URBANI SARS-CoV isolate was from a physician infected by a patient who contracted SARS while staying at Hotel M. Isolates CUHKW1, GZD1, BJ01, BJ02, BJ03, BJ04, however, came from patients with no known linkage with Hotel M and, on the whole, were derived later than the Hotel M linked set. Our results showed that the cases associated with exposure in the Hotel M formed a cluster that was distinct from the other isolates.

In addition to this four locus genotype, the variant sequences at position 19 084 distinguished the Singaporean isolates from all others. There also seems to be a signature for the North China isolates at positions 9854, 19 838, and 27 243. Thus, the common sequence polymorphisms might be useful in identifying the differential source of a SARS viral infection.

Whether any of these common polymorphisms will result in biological and clinical differences remains to be determined. However, the common mutation in position 22 222 changing an isoleucine residue to threonine in the important antigenic region of the spike protein might be relevant. Mutations in this region of the SARS-CoV genome can arise because of selective pressure from host immune responses. That an isoleucine is present in all Hotel M linked isolates whereas a threonine at the same position in the major antigenic protein is found in all other geographically distinct isolates suggests that such non-conservative aminoacid changes have occurred to evade immunological pressures.

The SARS viral epidemic has placed a substantial strain on the health and economic status of nations. Understanding the nature of this virus and deriving methods to control the epidemic are very important. Our results show several molecular facets of the SARS coronavirus pertinent to public-health management of this epidemic. Its novelty as a human pathogen suggests that most populations might be immunologically naive to its infection. The discovery of genotypes linked to geographic and temporal clusters of infectious contacts suggests that molecular signatures can be used to refine contact histories.

Contributors

A E Ling organised and directed the in-vitro investigations of the biology of the SARS virus. K P Chan provided viral cultures, and L L Oon and S Y SeThoe did the molecular diagnostic analysis and extraction of the viral nucleic acids. N M Lee did the electron microscopy of the initial viral samples. Y S Leo was the chief infectious disease clinician caring for the patients. Sequence determination and analysis was done by Y Ruan, C Wei, H Thoreau, P Ng, K Chiu, L Lim, T Zhang, E E C Ren, L Stanton, P Long, and E T Liu.

Conflict of interest statement
None declared.

Acknowledgments

We thank Prasanna Kolatkar, Marie Wong, Liu Jianjun, and Joshy George for their help in this project. The authors are members of the GIS and are therefore funded by the Agency for Science Technology and Research of Singapore.

References

- 1 WHO. Cumulative number of reported probable cases of severe acute respiratory syndrome (SARS). www.who.int/csr/sarscountry/2003_04_24/en/ (accessed April 19, 2003).
- 2 Drosten C, Gunther S and Preiser W, et al. Identification of a novel coronavirus in patients with severe acute respiratory syndrome. *N Engl J Med* 2003 (published online April 10, 10.1056/NEJMoa030781).
- 3 Ksiazek TG, Erdman D, Goldsmith CS, et al. A novel coronavirus associated with severe acute respiratory syndrome. *N Engl J Med* 2003; (published online April 10, DOI: 10.1056/NEJMoa030747).
- 4 Peiris JSM, Lai ST, Poon LLM, et al. Coronavirus as a possible cause of severe acute respiratory syndrome. *Lancet* 2003; **361**: 1319–25.
- 5 Kuo L, Godeke GJ, Raamsman MJ, Masters PS, Rottier PJ. Retargeting of coronavirus by substitution of the spike glycoprotein ectodomain: crossing the host cell species barrier. *J Virol* 2000; **74**: 1393–406.
- 6 Sanchez CM, Izeta A, Sanchez-Morgado JM, et al. Targeted recombination demonstrates that the spike gene of transmissible gastroenteritis coronavirus is a determinant of its enteric tropism and virulence. *J Virol* 1999; **73**: 7607–18.
- 7 Phillips JJ, Chua MM, Lavi E, Weiss SR. Pathogenesis of chimeric MHV4/MHV-A59 recombinant viruses: the murine coronavirus spike protein is a major determinant of neurovirulence. *J Virol* 1999; **73**: 7752–60.
- 8 Bergmann CC, Yao Q, Lin M, Stohlman SA. The JHM strain of mouse hepatitis virus induces a spike protein-specific Db-restricted cytotoxic T cell response. *J Gen Virol* 1996; **77**: 315–25.
- 9 Gomez N, Carrillo C, Salinas J, Parra F, Borca M V, Escribano JM. Expression of immunogenic glycoprotein S polypeptides from transmissible gastroenteritis coronavirus in transgenic plants. *Virology* 1998; **249**: 352–58.
- 10 Jackwood MW, Hilt DA. Production and immunogenicity of multiple antigenic peptide (MAP) constructs derived from the S1 glycoprotein of infectious bronchitis virus (IBV). *Adv Exp Med Biol* 1995; **308**: 213–19.
- 11 Ndifuna A, Waters A K, Zhou M, Collison E W. Recombinant nucleocapsid protein is potentially an inexpensive, effective serodiagnostic reagent for IBV. *J Virol Methods* 1998; **70**: 37–44.
- 12 Callenbaut P, Enjuanes L, Pensaert M. An adenovirus recombinant expression the spike glycoprotein of porcine respiratory coronavirus is immunogenic in swine. *J Gen Virol* 1998; **77**: 309–13.
- 13 Homberger FR. Nucleotide sequence comparison of the membrane protein genes of three enterotropic strains of mouse hepatitis virus. *Virus Res* 1994; **31**: 49–56.
- 14 WHO. Severe acute respiratory syndrome (SARS). *Wkly Epidemiol Rec* 2003; **78**: 81–88. <http://www.who.int/wer/pdf/2003/wer7812.pdf> (accessed April 27, 2003).
- 15 US CDC and WHO. Update: Outbreak of severe acute respiratory syndrome, worldwide. *MMWR Morb Mortal Wkly Rep* 2003; **52**: 269–72. (www.cdc.gov/mmwr/preview/mmwrhtml/mm5212a1.htm accessed March 28, 2003).
- 16 www.who.int/cdsdiagnostics (accessed April 5, 2003).
- 17 Michael Smith Genome Science Centre. www.bcgsc.bc.ca (accessed April 13, 2003).
- 18 Thompson JD, Higgins DG, Gibson TJ. CLUSTAL-W—Improving the sensitivity of progressive multiple sequence alignment through sequence weighting, position-specific gap penalties and weight matrix choice. *Nucleic Acids Res* 1994; **22**: 4673–80.
- 19 Gish W. WU-BLAST (1996–2003). <http://blast.wustl.edu> (accessed April 25, 2003).
- 20 Swofford D L. PAUP. Phylogenetic analysis using parsimony (and other methods). Version 4. Sunderland, Massachusetts: Sinauer Associates, 2003.
- 21 Marra MA, Jones SJM, Astell CR, et al. The genome sequence of the SARS-associated coronavirus. Published online May 1, 2003: <http://www.sciencemag.org/cgi/rapidpdf/1085953v1.pdf> (accessed May 7, 2003).
- 22 Rowe CL, Fleming JO, Nathan MJ, Sgro JY, Palmenberg AC, Baker SC. Generation of coronavirus spike deletion variants by high-frequency recombination at regions of predicted RNA secondary structure. *J Virol* 1997; **71**: 6183–90.
- 23 Lee CW, Jackwood MW. Origin and evolution of Georgia 98 (GA98), a new serotype of avian infectious bronchitis virus. *Virus Res* 2001; **80**: 33–39.
- 24 Bush RM, Smith CB, Cox NJ, Fitch WM. Effects of passage history and sampling bias on phylogenetic reconstruction of human influenza A evolution. *Proc Natl Acad Sci USA* 2000; **97**: 6974–80.

The significance of the correspondence between genomic data and links with hotel M

One way that the content of the SARS genomes sequenced to date aligns with the clinical data is as follows. After a multiple alignment of the sequences of 13 of the SARS genomes (sequence BJ04 was excluded from this analysis because it seems to be incomplete), in four different positions, if you segregate the samples on the basis of which nucleotides are present, the division is the same as when the samples are segregated on the basis of whether the patient had any link to the Hotel M. Here, we assess the significance of this finding.

This analysis concentrates on point mutations occurring on the 26 140 loci in which a nucleotide is present and definitively sequenced in all 13 genomes. In the rest of this document, we will restrict ourselves to these loci without further comment, as if they constituted all of the genomes.

Let us refer to the eight samples that are associated with hotel M as the M-samples, and the other five as the $\bar{M}$ -samples (the non-M samples).

The null hypothesis is that nucleotides in all 13 genomes are obtained through mutation from a single consensus sequence. Specifically, the null hypothesis is that each position in each genome is mutated independently at random with a common probability q , and each of the three possible nucleotides are equally likely.

The value of q used in the analysis was obtained as follows. We took, for the purpose of calculating q , a consensus sequence obtained by choosing the most commonly occurring nucleotide at each position. Then, we calculated the total, over all genomes, of the number of departures from the consensus: this number was 104. Then, in keeping with the null hypothesis we set $q=104 \div (13 \times 26\ 140)$.

Our first calculation will bound the probability p that there are at least four loci with genetic variation that aligns perfectly with (M, $\bar{M}$) segregation of the samples. Let us begin by evaluating the probability that a single locus will have this property—let us fix our attention on a single, arbitrary, locus L for this purpose for the moment. The ways that this could happen can be divided into three groups:

- The same mutation could happen in each of the M samples, with no mutations in any of the $\bar{M}$ samples,
- The same mutation could happen in each of the $\bar{M}$ samples, with no mutations in any of the M samples, or
- The same mutation could happen in each of the M samples, and the same mutation could happen in all of the $\bar{M}$ samples.

Each of the first two could be further divided on the basis of which of the non-consensus nucleotides are present in the M and $\bar{M}$ samples, respectively. In this way, we can see that the probability that the nucleotides at locus L segregate the samples in the same way as the (M, $\bar{M}$) distinction is at most

$$3(1-q)^{8\bar{M}_1}(q \div 3)^{8M_1} + 3(1-q)^{8M_1}(q \div 3)^{8\bar{M}_1} + 6(q \div 3)^{8\bar{M}_1+8M_1} = 3(1-q)^8(q \div 3)^8 + 3(1-q)^8(q \div 3)^8 + 6(q \div 3)^{13}$$

The null hypothesis includes the assertion that what happens at different loci are independent. Thus, we can apply the following large deviation bound.

Lemma 1 (Theorem A.12 of [1]) Suppose a biased coin with probability r of coming up heads is tossed independently at random m times. Then for any $\beta \geq 1$,

$$\Pr(\text{number of heads} \geq \beta(rm)) \leq (e^{\beta-1} \beta^{-\beta})^{rm}.$$

In this case, we are interested in the probability that at least four loci segregate as the $(M, \overline{M})$ distinction. Thus, the relevant value of β is

$$\frac{4}{rm} = \frac{4}{(3(1-q)^5(q \div 3)^8 + 3(1-q)^8(q \div 3)^5 + 6(q \div 3)^{13}) \times 26\,140}$$

Using the above Lemma with this value of β ,

$$r = \frac{3(1-q)^5(q \div 3)^8 + 3(1-q)^8(q \div 3)^5 + 6(q \div 3)^{13}}{26\,140}, \quad \text{and } m = 26\,140,$$

we get that the probability p of at least four loci segregating as $(M, \overline{M})$ satisfies,

$$p \leq 10^{-60}.$$

One potential concern about this analysis is its dependence on q , whose value was estimated from the data. However, the conclusion of significance is not sensitive to the choice of q . If, instead of the estimated value of $\frac{104}{13 \times 26\,140} \approx 0.0003$, we instead took $q = 0.003$, ten times larger,

$$p \leq 10^{-40}.$$

Another limitation of this analysis is that the null hypothesis includes the assertion that different loci are mutated independently. This can be removed by using Markov's inequality in place of Lemma 1. Markov's inequality states that, for any nonnegative random variable X , $\Pr(X > aE(X)) \leq 1/a$. Since, if the samples at each locus are mutated independently of one another, the above analysis implies that the expected number of loci that segregate as does the $(M, \overline{M})$ distinction is

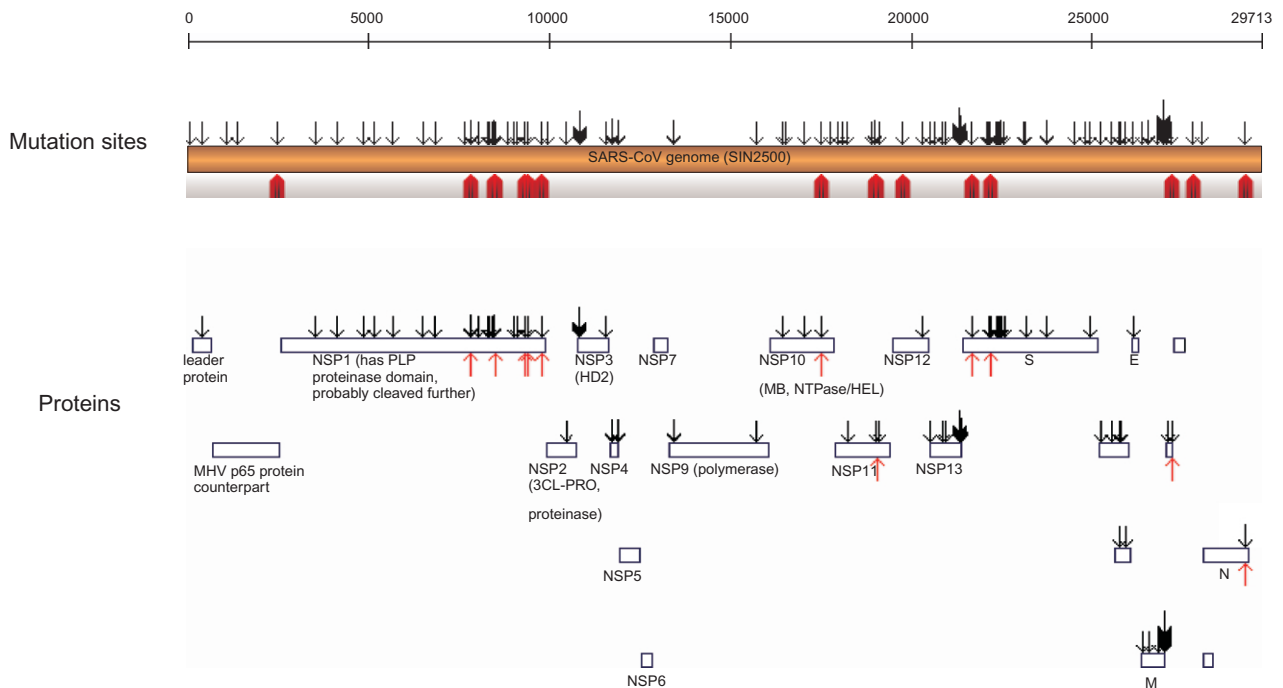
$$26\,140 \times (3(1-q)^5(q \div 3)^8 + 3(1-q)^8(q \div 3)^5 + 6(q \div 3)^{13}).$$

Thus, Markov's inequality implies that, under this weaker null hypothesis,

$$p < \frac{26\,140 \times (3(1-q)^5(q \div 3)^8 + 3(1-q)^8(q \div 3)^5 + 6(q \div 3)^{13})}{4} \leq 10^{-15}$$

References

- 1 N Alon, J H Spencer, and P Erdős. The Probabilistic Method. Wiley, 1992.



Sequence variation map

Sequence nucleotide variations occurred in the 14 SARS-CoV isolates were identified and located. Downward arrows (black) are for variations occurred in single isolate. Upward arrows (red) are for recurrent variations. Variations that cause aminoacid changes were also mapped on proteins.

Webappendix2: Genome sequence variations occurred in 14 SARS-CoV isolates

	Singapore cases					Overseas cases										variation frequency	ORF / protein	AA change
Genome position	Index SIN2500 29711	Primary SIN2748 29706	Primary SIN2774 29711	Primary SIN2677 29705	Secondary SIN2679 29711	Canada TOR2 29712	Hanoi URBANI 29727	HK HKU39849 29727	HK CUHKW1 29712	S China GZ01 29429	N China BJ01 28920	N China BJ02 29430	N China BJ03 29291	N China BJ04 24774				
17564	T	T	T	T	T	T	T	T	G	29429	G	G	G	G	6	nsp10	D→E	
27827	T	T	T	T	T	T	T	T	C	C	C	C	C	C	6	non-coding		
9404	T	T	T	T	T	T	T	T	C	C	C	C	C	-	5	nsp1	V→A	
22222	T	T	T	T	T	T	T	T	C	C	C	C	C	N	5	spike	I→T	
19084	T	T	T	T	C	C	C	C	C	C	C	C	C	C	4	nsp11	I→T	
19834	A	A	A	A	A	A	A	A	A	A	G	G	G	G	4	nsp12	silent	
9854	C	C	C	C	C	C	C	C	C	C	T	T	T	-	3	nsp1	A→V	
27243	C	C	C	C	C	C	C	C	C	C	T	T	N	T	3	putative protein	P→L	
2601	T	T	T	T	T	C	T	C	T	T	T	T	T	-	2	MHV p65	silent	
7919	C	C	C	C	C	C	T	C	C	C	C	C	T	C	2	nsp1	A→V	
8559	T	T	T	T	T	T	T	T	T	C	T	T	T	A	2	nsp1	Y→*	
8572	G	G	G	G	G	G	G	G	G	T	G	T	G	G	2	nsp1	V→L	
9479	T	T	T	T	T	T	T	T	C	C	T	T	T	-	2	nsp1	V→A	
19064	A	A	A	A	A	A	A	G	A	A	A	A	A	A	2	nsp11	silent	
21721	G	G	G	G	G	G	G	G	A	-	-	A	-	-	2	spike	G→D	
29279	A	A	A	A	A	A	A	A	A	C	A	C	A	A	2	Nucleocapsid	Q→P	
508	G	G	G	G	G	G	G	G	G	T	G	G	G	G	1	leader	G→C	
1206	T	T	T	T	T	T	T	T	T	C	T	T	T	-	1	MHV p65	silent	
1476	A	A	A	G	A	A	A	A	A	A	A	A	A	A	1	MHV p65	silent	
3626	T	T	T	T	T	T	T	T	T	C	T	T	T	T	1	nsp1	I→T	
4220	A	A	A	A	A	A	A	A	A	G	A	A	A	A	1	nsp1	K→R	
4952	T	T	T	T	T	T	T	T	T	T	T	A	T	-	1	nsp1	M→K	
5251	C	C	C	C	C	C	C	C	C	A	C	C	C	-	1	nsp1	L→I	
5773	A	A	A	A	A	A	A	A	A	N	A	C	A	-	1	nsp1	K→Q	
6612	G	G	G	G	G	G	G	G	G	T	G	-	G	N	1	nsp1	L→F	
6929	G	G	G	G	G	G	G	G	G	A	G	G	G	G	1	nsp1	C→Y	
7746	G	G	G	G	G	G	G	G	T	G	G	G	G	G	1	nsp1	silent	
7930	G	G	G	G	G	G	G	G	A	G	G	G	G	G	1	nsp1	D→N	
8124	T	T	T	T	T	T	T	T	T	T	T	T	A	T	1	nsp1	D→E	
8387	G	G	G	G	G	G	G	G	C	G	G	G	G	G	1	nsp1	S→T	
8417	G	G	G	G	G	G	G	G	C	G	G	G	G	G	1	nsp1	R→T	
8502	T	T	T	T	T	T	T	T	T	G	T	T	T	T	1	nsp1	C→W	
8580	A	A	A	A	A	A	A	A	A	T	A	A	A	A	1	nsp1	silent	
8946	T	T	T	T	T	T	T	T	T	A	T	T	T	N	1	nsp1	silent	
9095	C	C	C	C	C	C	C	C	C	T	C	C	C	C	1	nsp1	T→I	
9176	T	T	T	T	T	T	T	T	T	T	T	T	T	T	1	nsp1	V→A	
10029	G	G	G	G	G	G	G	G	G	A	G	G	G	-	1	nsp2	silent	
10550	A	A	A	A	A	A	A	A	A	A	A	C	A	-	1	nsp2	Q→P	
10921	G	G	G	G	G	G	G	G	G	A	G	G	N	-	1	nsp3	V→K	
10922	T	T	T	T	T	T	T	T	T	A	T	T	N	-	1	nsp3	V→K	
10923	T	T	T	T	T	T	T	T	T	G	T	T	N	-	1	nsp3	V→K	
10926	G	G	G	G	G	G	G	G	G	A	G	G	N	-	1	nsp3	G→I	
10927	G	G	G	G	G	G	G	G	G	A	G	G	N	-	1	nsp3	G→I	
10928	G	G	G	G	G	G	G	G	G	T	G	G	N	-	1	nsp3	G→I	
10929	C	C	C	C	C	C	C	C	C	T	C	C	C	-	1	nsp3	silent	
10930	A	A	A	A	A	A	A	A	A	G	A	A	A	-	1	nsp3	T→A	
11664	G	G	G	G	G	G	G	G	G	N	N	G	A	N	1	nsp3	M→I	
11811	G	G	G	G	G	G	G	G	G	G	-	G	A	-	1	nsp4	silent	
11815	T	T	T	T	T	T	T	T	T	T	-	T	G	-	1	nsp4	S→A	
11818	G	G	G	G	G	G	G	G	G	G	-	G	A	-	1	nsp4	V→I	
11971	G	G	G	G	G	G	G	G	G	G	N	A	G	G	1	nsp4	D→N	
11974	A	A	A	A	A	A	A	A	A	A	N	T	A	A	1	nsp4	I→F	
13494	G	G	G	G	G	G	G	A	G	G	G	G	G	G	1	nsp9 RdRp	V→S	
13495	T	T	T	T	T	T	T	T	G	T	T	T	T	T	1	nsp9 RdRp	V→S	
15807	T	T	T	T	T	T	T	T	T	T	T	T	T	G	1	nsp9 RdRp	C→G	
16521	G	G	G	G	G	G	G	G	G	A	G	G	G	G	1	nsp10	D→N	
16622	C	C	C	C	C	C	C	T	C	C	C	C	C	C	1	nsp10	silent	
17131	T	T	T	T	T	T	T	T	T	C	T	T	T	T	1	nsp10	L→S	
17846	C	C	C	C	C	C	C	C	C	T	C	C	C	-	1	nsp10	silent	
18065	G	G	G	G	G	G	G	A	G	G	G	G	G	-	1	nsp11	silent	
18176	C	C	C	C	C	C	C	C	C	T	C	C	C	C	1	nsp11	silent	
18282	C	C	C	C	C	A	C	C	C	C	C	C	C	C	1	nsp11	L→I	
18965	T	T	A	T	T	T	T	T	T	T	T	T	T	T	1	nsp11	silent	
19183	T	T	T	T	T	T	C	T	T	T	T	T	T	T	1	nsp11	V→A	
20363	G	G	G	G	G	G	G	G	G	G	G	G	T	G	1	nsp12	M→I	

Genome position	Singapore cases					Oversea cases									variation frequency	ORF / protein	AA change
	Index SIN2500 29711	Primary SIN2748 29706	Primary SIN2774 29711	Primary SIN2677 29705	Secondary SIN2679 29711	Canada TOR2 29712	Hanoi URBANI 29727	HK HKU39849 29727	HK CUHKW1 29712	S China GZ01 29429	N China BJ01 28920	N China BJ02 29430	N China BJ03 29291	N China BJ04 24774			
20596	A	A	A	A	A	G	A	A	A	A	A	A	A	N	1	nsp13	Q→R
20699	A	A	A	A	A	A	A	A	A	A	N	C	A	-	1	nsp13	silent
20704	G	N	G	G	G	G	G	G	G	G	N	G	G	-	1	nsp13	G→E
20910	G	G	G	G	G	T	G	G	G	G	G	G	G	-	1	nsp13	D→Y
20992	G	G	G	G	G	G	G	G	G	A	G	G	G	G	1	nsp13	R→K
21403	A	A	A	A	A	A	A	A	A	N	N	A	N	T	1	nsp13	silent
21404	T	T	T	T	T	T	T	T	T	N	N	T	N	C	1	nsp13	silent
21405	T	T	T	T	T	T	T	T	T	N	N	T	N	A	1	nsp13	S→K
21406	C	C	C	C	C	C	C	C	C	N	N	C	N	A	1	nsp13	S→K
21407	T	T	T	T	T	T	T	T	T	N	N	T	N	A	1	nsp13	S→K
21408	C	C	C	C	C	C	C	C	C	N	N	C	N	T	1	nsp13	L→C
21409	T	T	T	T	T	T	T	T	T	N	N	T	N	G	1	nsp13	L→C
21411	C	C	C	C	C	C	C	C	C	N	N	C	N	T	1	nsp13	L→C
21412	T	T	T	T	T	T	T	T	T	N	N	T	N	G	1	nsp13	L→C
21413	G	G	G	G	G	G	G	G	G	N	N	G	N	T	1	nsp13	L→C
21414	G	G	G	G	G	G	G	G	G	N	N	G	N	C	1	nsp13	E→P
21415	A	A	A	A	A	A	A	A	A	N	N	A	N	C	1	nsp13	E→P
21416	A	A	A	A	A	A	A	A	A	N	N	A	N	G	1	nsp13	E→P
21417	A	A	A	A	A	A	A	A	A	N	N	A	N	T	1	nsp13	silent
21418	A	A	A	A	A	A	A	A	A	N	N	A	N	G	1	nsp13	silent
22145	T	T	T	T	T	T	T	T	T	C	T	T	T	N	1	spike	silent
22207	C	C	C	C	C	C	C	C	C	T	C	C	C	N	1	spike	S→L
22422	G	G	G	G	G	G	G	G	G	A	G	G	G	G	1	spike	G→R
22466	T	T	T	T	T	T	T	T	T	T	N	G	T	T	1	spike	F→L
22479	A	A	A	A	A	A	A	A	A	A	N	A	T	A	1	spike	N→Y
22517	A	A	A	A	A	A	A	A	A	G	N	A	A	A	1	spike	silent
22522	A	A	A	A	A	A	A	A	A	G	N	A	A	A	1	spike	K→R
22634	T	T	T	T	T	T	T	T	T	T	T	T	A	T	1	spike	N→K
23174	C	C	C	C	T	C	C	C	C	C	C	C	C	C	1	spike	silent
23220	T	T	T	T	T	G	T	T	T	T	T	T	T	T	1	spike	S→A
23792	C	C	T	C	C	C	C	C	C	C	C	C	C	N	1	spike	silent
23823	T	T	T	T	T	T	T	T	T	G	T	T	T	N	1	spike	Y→D
24566	T	T	T	T	T	T	T	T	T	C	T	T	T	T	1	spike	silent
24872	T	T	T	T	T	T	C	T	T	T	T	T	T	T	1	spike	silent
24978	A	A	A	A	A	A	A	A	A	G	A	A	A	A	1	spike	K→E
25299	G	G	G	G	G	G	G	G	G	G	G	A	G	G	1	putative protein	G→E
25569	T	T	T	T	T	T	T	A	T	T	T	T	-	T	1	putative protein	M→K
25779	A	A	A	A	A	A	A	A	A	C	A	A	A	A	1	putative protein	E→A
25844	A	A	A	A	A	A	A	A	A	T	A	A	A	A	1	putative protein	R→W
25984	C	C	C	C	C	C	C	C	C	C	C	T	C	C	1	putative protein	silent
26186	G	G	G	G	G	G	G	G	G	A	G	G	G	G	1	protein E	V→M
26428	A	G	G	G	G	G	G	G	G	G	G	G	G	G	1	protein M	K→E
26586	T	T	T	T	T	T	T	T	T	C	T	T	T	T	1	protein M	silent
26600	C	C	C	C	C	C	C	T	C	C	C	C	C	C	1	protein M	silent
26857	T	T	T	T	T	T	C	T	T	T	T	T	T	T	1	protein M	S→P
27012	A	A	A	A	A	A	A	A	A	N	A	A	A	G	1	protein M	silent
27013	A	A	A	A	A	A	A	A	A	N	A	A	A	T	1	protein M	N→Y
27016	A	A	A	A	A	A	A	A	A	N	A	A	A	G	1	protein M	T→G
27017	C	C	C	C	C	C	C	C	C	N	C	C	C	G	1	protein M	T→G
27019	G	G	G	G	G	G	G	G	G	N	G	G	G	A	1	protein M	D→N
27022	C	C	C	C	C	C	C	C	C	N	C	C	C	T	1	protein M	H→Y
27024	C	C	C	C	C	C	C	C	C	N	C	C	C	T	1	protein M	H→Y
27025	G	G	G	G	G	G	G	G	G	N	G	G	G	A	1	protein M	A→K
27026	C	C	C	C	C	C	C	C	C	N	C	C	C	A	1	protein M	A→K
27027	C	C	C	C	C	C	C	C	C	N	C	C	C	A	1	protein M	A→K
27028	G	G	G	G	G	G	G	G	G	N	G	G	G	T	1	protein M	G→L
27029	G	G	G	G	G	G	G	G	G	N	G	G	G	T	1	protein M	G→L
27030	T	T	T	T	T	T	T	T	T	N	T	T	T	A	1	protein M	G→L
27032	G	G	G	G	G	G	G	G	G	N	G	G	G	A	1	protein M	S→N
27033	C	C	C	C	C	C	C	C	C	N	C	C	C	T	1	protein M	S→N
27111	A	A	A	G	A	A	A	A	A	A	A	A	A	A	1	putative protein	E→G
28089	C	C	C	C	C	C	C	C	C	T	C	C	C	N	1	non-coding	
variations	2	2	3	3	2	4	6	9	9	46	7	19	16	35	163		
		deletion 27794-27798		deletion 27760-27765													

Note: The first 16 are recurrent variations. The rest occurred only in single isolates

NP_828853.1 K→Q
NP_828853.1 silent
NP_828853.1 T→I