

細菌ゲノム配列情報から施設内遺伝子情報基盤データベース の構築とその活用に関する調査研究

綿引 正則 内田 薫 磯部 順子 木全 恵子
金谷 潤一 加藤 智子 大石 和徳

Construction and Utilization of In-house Genetic Infrastructure Databases from Bacterial Genome Sequence Information

Masanori WATAHIKI, Kaoru UCHIDA, Junko ISOBE, Keiko KIMATA,
Jun-ichi KANATANI, Tomoko KATO, and Kazunori OISHI

要 旨 次世代シーケンサー (NGS) による塩基配列決定技術の世界的な普及に伴い、公的 DNA データベースに登録されるゲノム配列情報も増大している。NGS から生産される膨大な配列情報を処理するゲノム解析は、主に Linux コンピュータを用いて世界的に進展しており、バイオインフォマティクスの知識が必要である。地方衛生研究所 (地研) の細菌検査部門においても、病原細菌のゲノム情報の活用は、感染源調査や病原性に関する解析等、公衆衛生上有用な情報が得られると期待されており、近年、導入が進んでいる。一方で、ヒューマンインターフェイスに優れた Windows 環境は広く普及しているが、バイオインフォマティクスの知識と Linux コンピュータの使用がメインになるため、地研におけるゲノム解析の利用は限定的であった。しかし、Windows10 の登場により WSL(Windows Subsystem for Linux) 機能が追加され、Linux 用のアプリケーションが Windows 環境でシームレスに使用出来るようになり、地研におけるゲノム解析のためのバイオインフォマティクス技術が利用しやすくなった。Windows 環境で、インターネット上に存在する膨大なゲノム配列情報を地研の細菌検査部門で利用するためには、バイオインフォマティクスの知識が必要であり、普及の壁になっていると思われる。そこで、本研究では、膨大なゲノム情報の中から、地研の細菌検査部門で有効利用できる情報を抽出し、バイオインフォマティクスの知識に依存しなくても利用を可能とする in-house データベースの構築のためのワークフローの作成とその活用法について検討した。カルバペネム耐性腸内細菌科細菌検査、腸炎ビブリオ分離株の SNP 情報のデータベース化、およびこれまで蓄積してきた病原体検出情報も同じくデータベース化し、その活用法を検討したところ、検査結果の迅速化や的確な情報発信などに寄与することが明らかになった。まだまだ発展中の分野であり、今回示したワークフローの検証を含めて、日々の業務を支援するデータベースとしてどのように維持するかが今後重要である。

近年の塩基配列決定技術の進歩は著しく、DNA シーケンサと呼ばれる塩基配列を決定する装置が普及し、地方衛生研究所 (地研) においても、検査機器の一つとして普及している。塩基配列情報は、アデニン (A)、グアニン (G)、シトシン (C) およびチミン (T) を示す AGCT の4つの塩基文字列によるテキストデータとして表現されている。そのためデータベース化しやすく、日本をはじめ、米国や欧州の3極協調による塩基配列データベース (「公的 DNA データベース」と称される；日本の場合は、DDBJ[1] が担当) が構築されており、インターネット上で誰でも自由に利用することができるのが大きな特徴である。

公衆衛生領域に限れば、感染症の原因となる微生物の塩基配列情報も日々増加しており、その活用も広がっている。例えば、1) 細菌、真菌及び原虫のリボゾーム RNA 配列解析による分子同定、2) ウイルスや細菌遺伝子の配列多型解析による病原体の遺伝系統の解析や分子疫学情報として利用した感染源や感染ルート調査への応用、3) 公開された遺伝子配列から病原微生物を解析する PCR プライマーの設計等、既に日常的な検査に利用されている。

公的データベース上の配列情報は誰でも自由に利用することが出来る、有用性の高い情報である。それをどのように利用するか、特に近年増加する

病原微生物ゲノム情報は、細菌検査の現場で有用と思われるが、食中毒細菌（例えば、病原性大腸菌やサルモネラ属菌）のゲノムサイズは約4～6Mbp（メガ塩基対）と大きく、塩基配列を解析するコンピュータとソフトウェアが必須である。さらに近年、地研においても次世代シーケンサー（NGS）の導入が進み、バイオインフォマティクス技術の必要性が増大している。NGSから生産された大量の塩基配列データは、コンピュータのOSとして、WindowsではなくUnix系のLinux上で動く様々なアプリケーションが開発されゲノム解析に広く利用されているが、WindowsユーザーにとってLinuxの利用はこれまで難しかった。しかし、Linux用のソフトをそのまま使えるようにするWSL（Windows Subsystem for Linux）機能が利用できるようになったWindows10の普及により、NGSによる細菌ゲノム解析はWindows環境で利用できるようになってきている。

地研におけるゲノム解析は、地域で分離された病原ウイルスや病原細菌が対象となる。ウイルスと細菌のゲノムサイズは、大きく異なり、そのため、ゲノム解析の手法が異なる。本調査研究においては、地研における細菌検査部門に限定する。地研の細菌検査の現場で生産される日々の遺伝子検査結果と公的DNAデータベースから得られる膨大な塩基配列情報を利用するためには、その情報を如何に利用するかが重要である。適切なソフトウェアを用いても、かなりの時間を有する。従って、日々、検査の現場で利用できる有用な遺伝子情報を抽出し、施設内で簡便に利用できるデータベースの構築は有用と思われる。あらかじめ解析に利用できるin-house（施設内）データベースとして構築しておけば、必要な情報が短時間に得られ、結果的に検査時間の短縮につながると考えられる。

インターネット上で得られる情報は膨大であり、有用な情報が存在するものの、様々なファイル形式で存在しており、有効に利用するためには必要な情報のみを抽出することが必要となる。特に遺伝子に関する情報については、それぞれ目的に応じた拡張子や様々なファイル形式が定義されている。これらのファイル形式は基本的には文字列のみのテキストデータ（plain textとも呼ばれる）から成る。従って、このような文字列から、検査や調査研究に利用できる情報を効率的に抽出することが必要である。

本調査研究は、地研の検査の現場でコンピュー

タのOSとして広く利用されているWindows環境で、細菌検査や調査研究に有用となる病原体細菌の遺伝子情報に関するin-houseデータベースを構築する方法を検討した。さらに構築したデータベースの活用法についても検討を加えたので報告する。

材料と方法

1. 使用コンピュータの環境、Webツールおよびアプリケーション

1) 使用したコンピュータは、64-bit仕様のWindows 10で、Ubuntu on Windows Subsystem for Linux (WSL)の使用及びインターネット接続環境とした。

2) Web ツール

- ・DDBJ [1], NCBI [2]: 公的DNAデータベースで、あらゆる生物の解析されたDNA配列の登録、検索（blast 検索サービス）等の利用がインターネット経由で可能。登録DNA塩基配列はgenebankファイル（拡張子gbkあるいはgb）として、遺伝子の情報と配列が記載された書式（gbkファイル）としてダウンロード可能。
- ・RAST [3]: ゲノム解析された塩基配列に遺伝子と機能を自動的に割当て（annotation）、gbkファイルやExcelファイル形式でダウンロード可能。

3) 使用ソフト

① Windows環境で使用したソフトウェアおよび使用目的は以下の通りである。

- ・CLC Genomic Workbench（フィルジェン株式会社）:de novo アセンブル

- ・秀丸エディタ [4]: テキストデータの加工（正規表現、検索、置換、grep 検索）、csvファイルからfastaファイル*変換、文字コード・改行コードの変換（WindowsからLinux）
*fastaファイルとは、拡張子fasta, fasあるいはfa: 配列の識別名称とDNA配列を含む文字列より成るplain textのこと。

- ・エクセル（Microsoft Corp）: テキストデータの加工、抽出（関数使用）

- ・FileMaker（FileMaker Pro 16 Advanced; FileMaker, Inc）: 検索、関数を使用したテキストデータの加工とゲノム情報データベースを構築。

- ・Artemis [5]: gbkファイルを読み込み、CDS（coding sequence、蛋白質をコードする遺伝

子) 情報を視覚化およびCDS 配列 (塩基配列, アミノ酸配列) の抽出, gbk ファイルの編集.

- ・ Sequence Assistant ver2.7 [6]: 塩基配列の連結, DNA 配列の相補鎖配列の変換.

- ・ Ruby script[7]: 複数の塩基配列情報を含む multi-fasta ファイルの csv 形式への一括変換 (Windows 付属のコマンドプロンプト使用)

- ・ FastQC [8]: NGS データのクオリティチェック

- ・ trimmomatic [9]: Miseq で得られた配列情報からのアダプター除去と及び配列のトリミング.

② Linux (ubuntu on WSL) で使用したソフトおよび使用目的は以下の通りである.

- ・ nucmer [10]: 相同 (共通) CDS の抽出, および一塩基多型 (Single Nucleotide Polymorphism; SNP) の抽出

- ・ raxml-ng [11]: 系統樹作成

尚, 使用アプリケーションについては, Windows10, WSL 及び Web 上のサービスを利用した場合には, アプリケーション名や Web 上のサービス名の後に, 必要に応じて, それぞれ「/Win」, 「/Lin もしくは/Linux」 および「/Web」と付加して表記した.

2. 供試菌

SNP 解析, SNP データベースの構築とその評価のため, 腸炎ビブリオの保存株を使用した. 内訳は, 1997-8 年分離 4 株 (臨床株: 血清型 O4:K68 1 株, O3:K6 3 株), 2010 年分離 1 株 (海水からの分離株: 血清型 O3:K6) および 2016-7 年分離 3 株 (臨床株: 血清型 O3:K6), 合計 8 株である.

3. ゲノム解析, SNP 抽出

細菌分離株のゲノム配列は NGS である MiSeq (illumina, Inc) を使用し取得した. DNA ライブラリーは Nextera XT DNA Library Preparation Kit (illumina, Inc) を用いて作製し, 試薬キットは MiSeq Reagent Kit v3 (PE, 600 サイクル) を使用した. MiSeq から得られた fastq ファイルからゲノム配列の作成は, 品質トリミング後 (FastQC/Win, trimmomatic/Win), CLC Genomics Workbench/Win の De Novo アセンブルツールを用いて, contig 配列より成る fasta ファイルとして得た. 得られた配列から, 比較したい株すべてに共通に存在する CDS の抽出, および参照配列と比較した SNP の抽出には nucmer/Lin を用いた. 参照配列とゲノム配列) に共通す

る CDS を選択し, 共通の CDS のみからなる SNP 抽出用の参照配列を作成 (multi-fasta 形式) した. これを参照配列として, nucmer/Lin を利用して, 各ゲノム配列と比較し塩基の変異データを抽出し, 塩基置換部位を記述するファイル形式である vcf (variant call format) ファイル (参照配列に対する位置と塩基情報と比較配列の代替塩基情報等を含む) として出力した. この vcf ファイルから, 秀丸エディタを使って, CDS 名, 参照配列に対する位置, 参照配列塩基, 解析株の置換塩基を抽出し, csv ファイルに変換し, FileMaker にインポートして, SNP データベースとした. 今回は, 塩基置換だけでなく, 挿入と欠失 (Indel) 情報も得た.

本研究では, 腸炎ビブリオの保存株を用いた SNP データベースの構築と評価を行うため, 参照配列は *Vibrio parahaemolyticus* RIMD 2210633 のゲノム配列 (Accession No. BA000031, BA000032) を用いた.

4. 系統樹作成

FileMaker による SNP データベース上で, indel を排除し, すべての SNP 位置の塩基を Excel ファイルとしてエクスポートした. 参照配列と各株の SNP 塩基はそれぞれ Excel の連結関数を用いて連結し, 秀丸エディタにコピー, 貼り付けし, multi-fasta 形式のファイルを作成した. この fasta ファイルを用いて, raxml-ng/Lin を用いて系統樹を作成した. 系統樹の表示は, FigTree/Win[12] を用いた.

結 果

1. 遺伝子情報基盤データベースの構築: β - ラクタマーゼ遺伝子

FileMaker/Win を用いて, in-house データベースを構築した. 本データベースで収集する情報は, 主に蛋白質をコードする遺伝子 (CDS) であり, 付随する情報として, 菌種名, 株情報, accession 番号 (公的データベースで付与された受入番号), 遺伝子名, 塩基配列とアミノ酸配列及びそのサイズ (塩基数とアミノ酸数) を含み, その他, データベースの目的に応じて収集すべき情報を追加した. ここで主な情報元は, 3 つで, 1 つ目は, 公的データベース (DDBJ と NCBI) に登録されている gbk ファイル, 2 つ目は特定の遺伝子あるいは菌種に特化した遺伝子データベースで通常 fasta ファイルとして入手できる. 3 つ目

A. gbk-formated or fasta-formated files from the public database, convert to gbk file to csv file for extracting CDS data sets.

1. Viewing gbk files by Artemis/Win, extract CDS data and export to csv file.
2. Re-annotation of gbk files by RAST/Web, download Excel-formated file.
3. Download of fasta-formated specific gene databases and convert to csv file by Ruby script.

B. fasta-files of laboratory sequence data produced by NGS.

1. Annotation of gbk files by RAST/Web, download Excel-formated file.

Fig.1. Workflow of Construction of In-house Database using FileMaker/Win.

は、自施設で解析した配列データで、ゲノム解析の結果として、de novo アセンブル後、繋がった contig 配列の fasta ファイルとして得られる。登録されている細菌の完全ゲノム配列の gbk ファイルは、ファイル容量は通常約 10M バイトで、テキストデータとしては大きなファイルである。ここでこのようなファイルから抽出した情報を FileMaker にインポートすることは直接には不可能であり、FileMaker にデータインポートを可能とするファイル形式である、csv あるいは Excel ファイル形式に先ず変換が必要である。そこで本研究では、3つの情報元から得た配列情報を FileMaker にインポートするためのワークフローを Fig.1 に示した。このワークフローを利用して、薬剤耐性菌の遺伝子検査に有用と思われる β -ラクタマーゼ遺伝子情報データベースを構築した。 β -ラクタマーゼ CDS の塩基配列及びアミノ酸配列は、NCBI の Web サイト [13] からクラス A から D に分類されている β -ラクタマーゼの遺伝子配列およびアミノ酸配列の fasta 形式のファイルをそれぞれダウンロードし、ワークフローの工程に従い FileMaker にインポートし、 β -ラクタマーゼ遺伝子情報の in-house データベースとした (Fig.1A3)。 β -ラクタマーゼ遺伝子情報については、日々、新しい遺伝子情報が公的データベースに登録されており、データベースの更新や訂正を以下のように行った。DDBJ の検索 Web ツールである ARSA[1]を利用して、キーワード検索して新しい遺伝子情報を調査し、必要であれば gbk ファイルとしてダウンロードし、塩基配列とアミノ酸配列情報を抽出し、データベースに追記した (Fig.1 A1)。その結果、2,338 配列よりなる β -ラクタマーゼ遺伝子情報データベースを作成

した。その内訳は、 β -ラクタマーゼの機能別の分類では、carbapenemase 632 遺伝子、Extended Spectrum β -lactamase (ESBL) 488 遺伝子、AmpC β -lactamase 606 遺伝子およびその他の β -ラクタマーゼとして 662 遺伝子（公的データベースから得た gbk ファイル中の情報から特定できなかった β -ラクタマーゼ）を含み、クラス別収録遺伝子の分類では、 β -ラクタマーゼとして、classA 型 845 遺伝子、classB 型 252 遺伝子、classC 型 612 遺伝子、classD 型 608 遺伝子 およびその他遺伝子（公的データベースから得た gbk ファイル中の情報から特定できなかった β -ラクタマーゼ）から構成されている in-house データベースとなった。

2. β -ラクタマーゼ遺伝子配列データベースの活用

このデータベースを用いて、PCR に使用するプライマー配列が CDS 中に存在するかどうか確認が可能であり、日常的に実施される薬剤耐性菌の β -ラクタマーゼ遺伝子の PCR 検査に活用が期待できる。また、このデータベースからエクスポートして得られる β -ラクタマーゼ遺伝子配列の fasta ファイルは、ローカル blast 検索で利用可能なデータベースファイルに変換すると、検査で得られた β -ラクタマーゼ遺伝子の塩基配列から、 β -ラクタマーゼ遺伝子の遺伝子型別を簡単に行うことが可能であった。

**3. Windows による SNP データベースの構築：
腸炎ビブリオ分離株のゲノム解析と SNP
データベース**

分離株の高精度分子疫学的マーカーとして利用可能なゲノム SNP データベースを Windows 環境の FileMaker で構築した (Fig.2)。構築法とデー

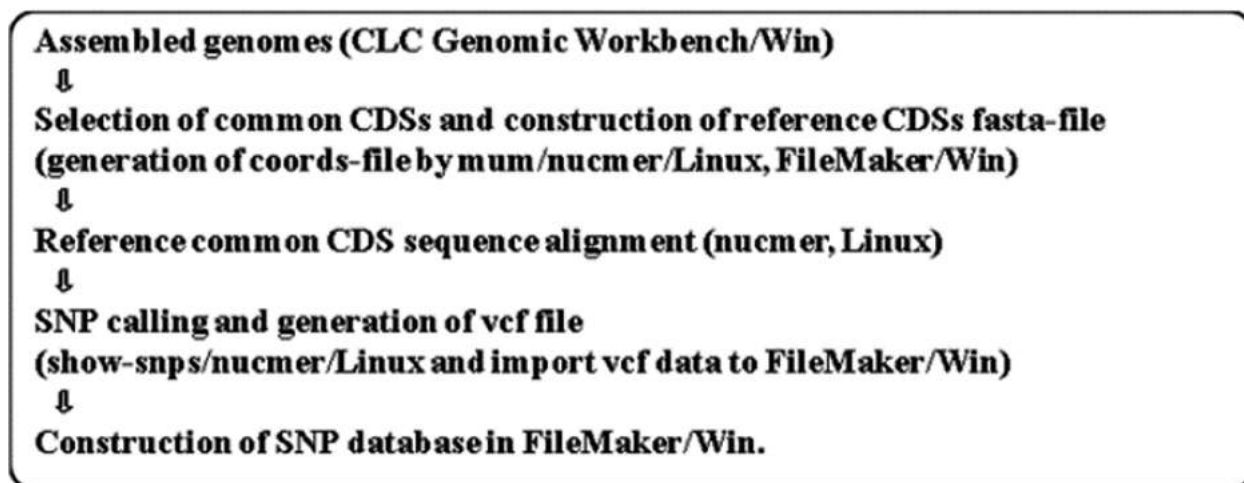


Fig.2. Workflow of Construction of SNP Databases using FileMaker/Win.

データベースの評価を行うため、1997-8年と2016-7年に分離された腸炎ビブリオ食中毒患者及び海水から分離された8株の腸炎ビブリオのゲノム配列をMiSeqより得た。得られたfastqファイルは、アダプター配列除去、品質評価後のトリミングを実施し、CLCgenomic Workbench/Winでde novoアセンブルを行い、40～60 contig配列を得た。得られた配列は、参照配列として、既に報告されているゲノム配列 (BA000031及びBA000032)と比較して、塩基置換部位 (single nucleotide variation site) を検索し、これをSNP (single nucleotide polymorphism) として利用した。SNPの抽出は、蛋白質をコードする領域 (CDS) を対象とし、2,254個のCDSからなる参照配列と8株のゲノム配列に共通に存在するCDSを対象として選択した。この共通するCDSの選択には、nucmer/Lin (-mum, オプションスイッチ使用) を使って、当所の供試8株のcontig配列から、今回は100%配列が一致したCDS (一塩基多型が存在しないため) を除き、また、2コピーあるいは部分的に2コピー以上存在するCDSも除いたところ、最終的に171個のCDSを得た。そこで参照配列からこの171個のCDS塩基配列を抽出し、multi-fastaファイルとして、これを今回、NGS解析した8株の腸炎ビブリオゲノム配列からSNP抽出に使用した。SNP抽出には、nucmerを用いてそれぞれの株の配列毎に、参照配列に対する塩基置換位置と置換塩基を検出し、vcfファイル形式で結果を得た。この結果を秀丸エディタとエクセルを用いて、必要な情報 (CDSの識別名称、CDS内の変異部位の位置情報、参照配列及び解

析株配列の塩基情報) を抽出し、FileMakerに逐次インポートし、SNPデータベースを構築した。複数株の塩基置換部位について、FileMakerへインポート時に参照配列と照合し、同一部位で既に検出されているか、あるいは新しい部位かを判定し、新しい場合にはデータを逐次追加するインポート法を採用した。今回のnucmerによる塩基変異部位の検出は、indelも検出しているため、FileMakerにインポートした後、このindelについてはすべて除去した。その結果、今回、171個のCDS (合計塩基数、180,648塩基) から、881か所でSNP座を検出した (Fig.3)。参照配列の塩基と、参照配列と較べてSNPとして検出された塩基は、大文字で、参照配列と同一塩基は小文字で標記した。また、ここで示した参照配列の「Ref_LocusCDS」フィールドで、VP239は771塩基のCDSであることを示しており、ここでは9か所のSNPが存在することを示されている。ここで例示したVP0239は *Vibrio parahaemolyticus* RIMD 2210633の第一染色体にコードされている triosephosphate isomeraseであった。

4. WindowsによるSNPデータベースの構築： 腸炎ビブリオ分離株のSNP配列による系統解析

SNPデータベースのデータを用いた系統解析のワークフローはFig.4に示した。参照株を含めたすべての株の881か所のSNPを含む塩基は、複数塩基配列をもつfasta形式のファイル (multi fasta) に変換した。作成したfastaファイルは、文字コード及び改行コードをLinux形式に変換し、raxml-ng/Linを用いて、系統樹を作成した。

Vp_SNP_db_EditedV1

ファイル(F) 編集(E) 表示(V) 挿入(I) 書式(M) レコード(R) スクリプト(S) ツール(T) ウィンドウ(W) ヘルプ(H)

23 881 合計 (ソート済み) レコード

すべてを表示 新規レコード レコード削除 検索 ソート 共有

レイアウト: VpSNP_Input original 表示方法の切り替え: プレビュー

Ref_LocusCDS	Pos	CDS_size	Ref	VP207	VP221	VP250	VP297	VP562	VP570	VP581	VP585	LocusCDS計算	LocusCDS集計
VP0003	132	1623	T	t	A	t	t	t	t	t	t	1	881
VP0015	76	114	A	a	a	a	a	a	a	T	a	1	881
VP0041	1293	2016	C	c	T	c	c	c	c	c	c	1	881
VP0083	25	198	C	c	A	c	c	c	c	c	c	1	881
VP0086	174	399	T	t	t	t	t	t	C	t	t	1	881
VP0147	431	726	A	a	a	a	a	a	G	G	a	1	881
VP0239	132	771	A	a	a	T	a	a	a	a	a	9	881
VP0239	270	771	A	a	a	T	a	a	a	a	a	9	881
VP0239	324	771	C	c	c	T	c	c	c	c	c	9	881
VP0239	366	771	T	t	t	C	t	t	t	t	t	9	881
VP0239	405	771	A	a	a	G	a	a	a	a	a	9	881
VP0239	567	771	T	t	t	C	t	t	t	t	t	9	881
VP0239	606	771	A	a	a	T	a	a	a	a	a	9	881
VP0239	615	771	A	a	a	G	a	a	a	a	a	9	881
VP0239	702	771	T	t	t	G	t	t	t	t	t	9	881
VP0240	73	348	G	g	g	T	g	g	g	g	g	7	881
VP0240	177	348	C	c	c	T	c	c	c	c	c	7	881
VP0240	243	348	G	g	g	A	g	g	g	g	g	7	881
VP0240	244	348	C	c	c	A	c	c	c	c	c	7	881
VP0240	257	348	T	t	t	C	t	t	t	t	t	7	881
VP0240	285	348	G	g	g	A	g	g	g	g	g	7	881
VP0240	307	348	C	c	c	A	c	c	c	c	c	7	881

Fig.3. Partial Depiction of SNP Database of *Vibrioparahaemolyticus* Isolates in Toyama
Dotted box shows the details of the SNP locus present in the reference CDS, VP0239.
See text for details.

Export SNP data to Excel format after searching on FileMaker/Win.
↓
Convert the concatenated SNP data sets to multi-fasta format file (Excel, Text editor/Win).
↓
Construction of phylogenetic trees (raxml-ng/Linux)
↓
Viewing of trees (FigTree/Win)

Fig.4. Construction of SNP phylogenetic tree

(Fig.5).

1997-8 年 の O3:K6 流 行 株 (VP207, VP221, VP297) と 2016-7 年 の O3:K6 株 (VP570, VP581, VP585) に加えて, 2010 年に海水から分離 O3:K3 株 (P562), 血清型が異なる VP250 (O4:K68) の 8 株では, 1997-8 年の分離株間の SNP は 6 ~ 13

違いで系統関係が極めて近いことが示されたが, 年代が上るとその SNP の違いは 100 前後であることが示された. 血清型が異なる分離株は, 遺伝的に離れており, 今回解析した血清型の異なる VP250 株と参照株 (Ref) の SNP 数は 658 であった. 今回, 解析した株のなかで最も参照配列

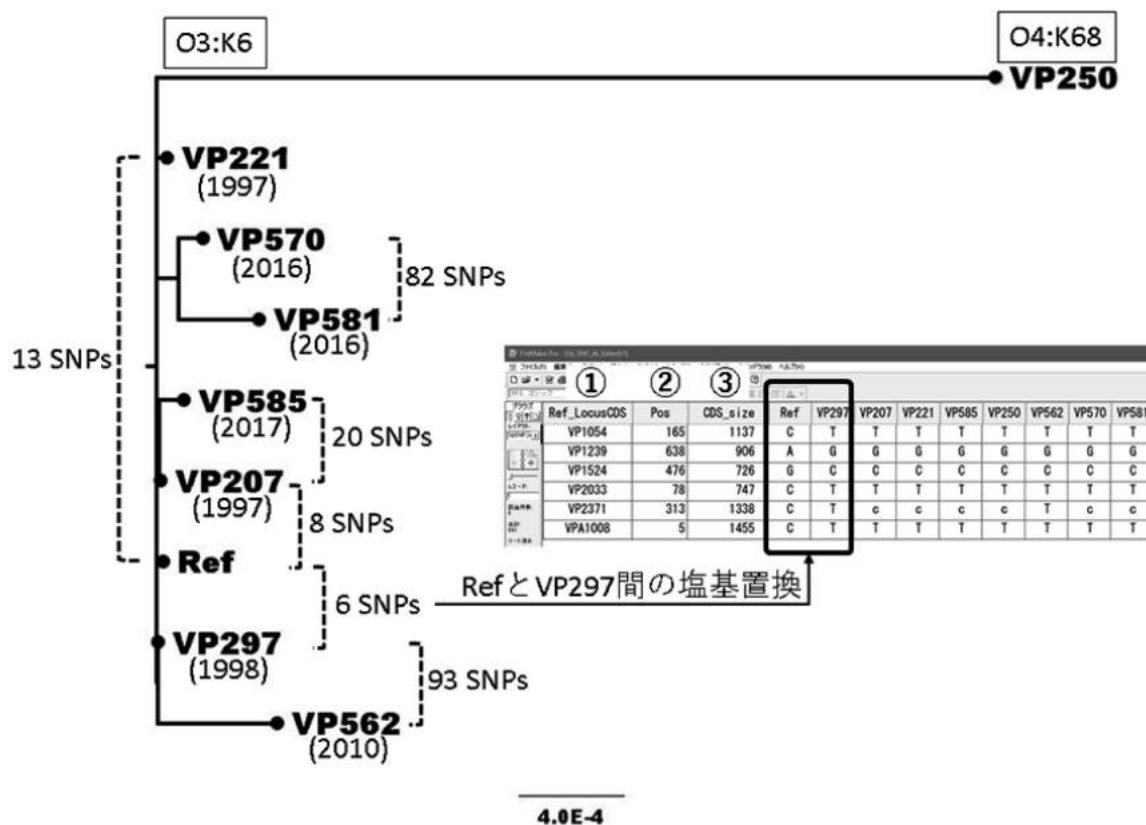


Fig.5. SNP-based Phylogenetic Tree of *Vibrio parahaemolyticus* Isolates in Toyama

Internal figure is described 6 SNPs after searching loci of base substitutions by FileMaker.

(Ref) に近縁であったのはVP297 株で、6 SNP の違いであり、PFGE パターンが同一になる可能性もあると推測される。本データベースから、参照配列のCDS における6つのSNP 位置の特定が可能であった (Fig.5)。

5. 病原微生物検出情報データベースの構築

検体別に病原細菌の検出数の集計が記録されている2000 年からの富山県病原微生物検出情報から、菌種名、分離年、分離月、分離材料に関する情報をFileMaker にインポートして、データベースを構築した。令和2 年3 月31 日時点で、報告された分離菌株総数は、350,711 株となった。検査材料別分離株数の割合は、尿及び呼吸器材料からの検出が多かった (Fig.6)。このデータベースを用いることにより、富山県で検出されている病原細菌の動態を迅速に把握することが可能となった。

考 察

細菌検査は、感染症や食中毒の原因となる病原細菌を分離し、分離した細菌の特徴を明らかにし、疾病診断や治療に資することが目的である。その中で、地研における細菌検査は公衆衛生上、感染症や食中毒が発生した場合、それが集団感染なのか、集団発生の場合、その感染源や感染ルートを明らかにし、蔓延防止に資するため、病原細菌の分離とその病原細菌の特徴を明らかにすることが求められる。また、平時においては、食品や環境中に存在する病原微生物のサーベイランスなどを行い、環境中のリスク評価を行なっている。その中で遺伝子検査は、その迅速性と再現性に優れ、得られる情報量の多さから近年、急速に普及が進み、塩基配列決定技術の需要は増大しつつある。また、ゲノム塩基配列情報が増大しているなかで、我々の検査や研究活動から生産される塩基配列情報の有効利用はいまだ限定的である。特にバイオインフォマティクスの知識が必要で、主にLinux

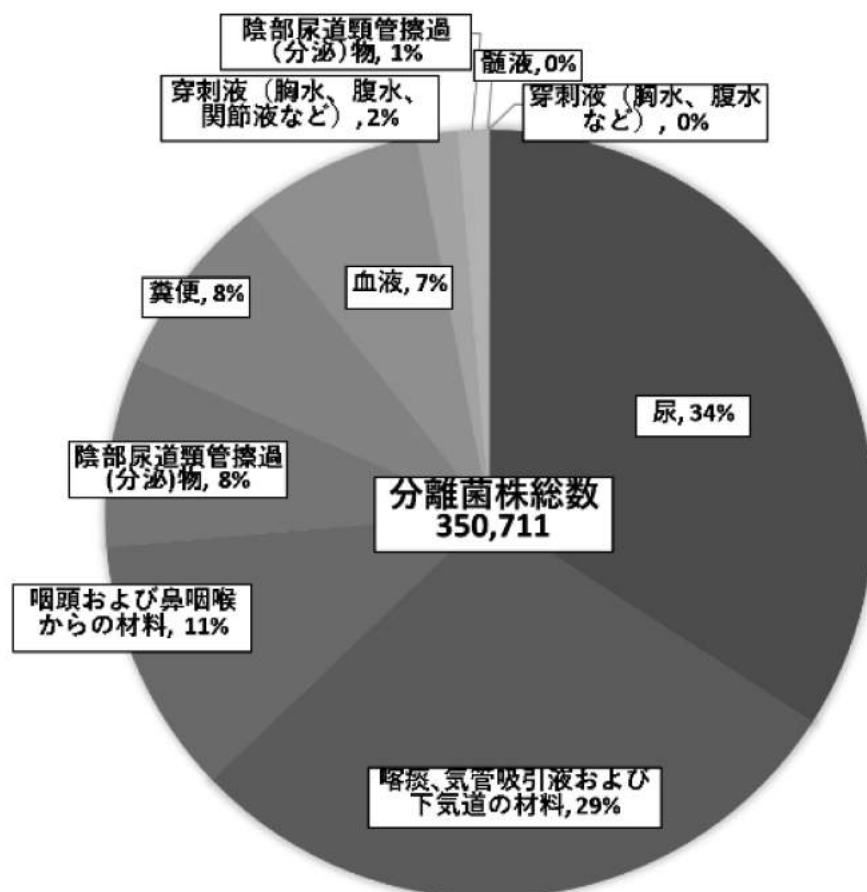


Fig.6. Isolation Rates of Pathogenic Bacteria by Clinical Specimen Detected in Toyama from 2000 to 2019.

環境で発展しており、Windowsを主に使っている現場ではその普及も限定的であった。しかし、最近発表されたWindows10の機能として追加されたWSLを利用すると、これまでLinux上で利用されてきたバイオインフォマティクス用のアプリケーションがWindowsとLinuxをシームレスで利用できるようになった。そこで本研究では、我々が必要とする情報をWindows環境で広く利用されているテキストエディタやExcelと言ったソフトの機能を使用し、様々なファイル形式で記述されている大量の情報から必要な情報を抽出し、我々が日常業務で簡単に利用できる情報に特化したデータベースの構築するためのワークフローの構築を試みた。また、その活用法を評価することも本研究の目的とした。

今回、作成したデータベースは、CRE検査に有用と思われる β -ラクタマーゼ遺伝子に関する情報、分離株の高精度分子疫学マーカーとして利用価値が高いといわれるSNP情報、さらに本県で長期にわたり蓄積されている病原微生物検出情報

に関する。何れも、FileMakerの検索機能が利用可能で、検索して得られた情報をエクスポートすることにより、blast検索や遺伝子配列に基づく遺伝子系統解析などに、すぐに利用することが可能になるというのが大きな利点である。

β -ラクタマーゼのデータベースは、カルバペネム耐性腸内細菌科細菌(CRE)感染症の起原因菌の検査に有用である。地研のCREの検査は、カルバペネマーゼ産生腸内細菌科細菌(CPE)を鑑別するため、スクリーニングとしてPCR[14]を実施し、カルバペネマーゼ遺伝子の有無を確認している。CPEと決定され、医療関連感染が疑われる場合には、カルバペネマーゼ遺伝子の亜型の決定が必要になる場合がある。今回、構築した β -ラクタマーゼのデータベースから、カルバペネマーゼ遺伝子配列と亜型情報を含めた遺伝子名を出力し、予め、ローカルブラスト検索のデータベースとして準備しておく、瞬時にカルバペネマーゼ遺伝子の鑑別が可能であった。この β -ラクタマーゼのデータベースは、地研の検査の現場

では有用であり、既に複数の地研でこのデータベースの情報が利用されている。また、このデータベースを利用して、新しいPCR 検査の開発、評価に利用された[15]。

地研で有用と思われる細菌のゲノム解析の一つに SNP 解析が挙げられる。地研の役割は、分子疫学情報として利用できる遺伝子型を決定することであり、アウトブレイクの早期探知、感染源や感染ルートの特定に欠かせない検査である。この目的に実施されている遺伝子検査法として、パルスフィールドゲル電気泳動 (PFGE) [16], IS-Printing System [17], Mutiple-Locus Variable number tandem repeat Analysis (MLVA) [18], Multi Locus Sequencing Typing (MLST) [19] などが知られている。しかし、菌種やその検査原理により、迅速性、使い勝手、解像度やコストなどが異なり、一長一短があることも事実である。特に、PCR 法をベースにした方法は、迅速性は優れているものの解像度が低い傾向にあり、ターゲットとする遺伝子により、分子疫学マーカーとして有用性が異なり、個別に評価が必要である。しかし、一方でコストやバイオインフォマティクスの知識が必要となるが、ゲノム配列上の一塩基多型を利用した解析として、高精度疫学マーカーとして利用できる SNP 解析が注目されている。これまで既存の解析法では解像度の点で問題のあったサルモネラ属菌や結核菌は、その有用性が報告されている [20, 21]。ここで SNP 抽出は、基準となるゲノム配列 (参照配列) に対して、NGS から得られた配列を張り付けて変異塩基を抽出する。また、

参照配列として、利用するゲノム配列により、core genome (cg) MLST, whole genome (wg) MLST [22] など SNP 解析には、操作性や使い勝手により様々な方法が報告されるようになっていく。そこで、今回、データベース化しやすいと思われる、共通の遺伝子 (CDS) から SNP を抽出する cgMLST の考え方をベースとした方法を開発した。その大きな特徴は、WSL の Linux 上で SNP 抽出を行い、得られた SNP データを逐次保存可能で、どの遺伝子上の SNP として検出されたかを知ることができる。今回使用した菌種は、腸炎ビブリオ食中毒で分離された研究室保存株である。本菌における、富山県の臨床分離株の検出株数は、1998 年に 252 件とピークで、当時世界的に流行し、公衆衛生上問題となった代表的な食中毒起因菌である (Fig.7)。当時から、PFGE 解析ではほとんど区別できない流行株として血清型 O3:K6 が注目されていた。この食中毒は、最近ではほとんど報告されなくなってきたが、2016 年に同じ血清型が 16 件検出された。約 20 年前に分離された同血清型の臨床分離株の遺伝的関係を調べるため、SNP 解析を実施すると同時に SNP データベースを構築し、その有用性を評価した。その結果 1997-8 年の分離株間の SNP は 6 ~ 13 違いで系統関係が極めて近いことが示されたが、年代が上がるとその SNP の違いは 100 前後であることが示された。また、このデータベースは、参照配列の CDS と塩基置換部位を簡単に知ることができるため、塩基置換の蓄積しやすい遺伝子など推定が可能になると期待される (Fig.5)。現在、本研究のワークフローを用いて、レジオネラ菌の SNP データベースを構築し、地域で分離された株間遺伝的関係を解析するツールとしての、本 SNP データベースの有用性の評価を進めている。

FileMaker によって構築した細菌ゲノム配列情報は、地研で実施されている遺伝子検査を支援し、迅速で、的確な検査結果を提供することが期待され、今後、他の菌種についても作成中である。また、我々は、これまで感染症サーベイランスの一環として、以前から県内の定点医療機関で分離された病原微生物の検出情報を集計し、情報還元し、NESID (感染症サーベイランスシステム) のサブシステムである病原体検出情報システムにも報告している。これらの情報も FileMaker に移行し、現在、効率的な情報還元方法を検討している。今後、当所のホームページを利用した情報還元等を利用していきたい。

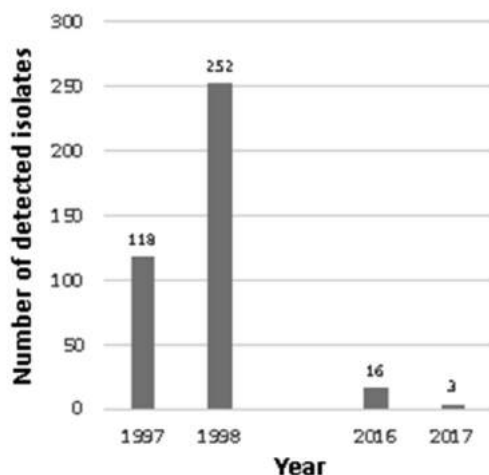


Fig.7. *Vibrioparahaemolyticus* Detected in Toyama at Specific Times.

今回、細菌ゲノム配列情報に基づいたin-houseデータベースの構築と活用法について報告した。細菌ゲノム配列情報は、日々、増大しており、新しい情報が追加されており、感染症や食中毒の流行する菌種や感染様式も時代とともに変化しているのも事実である。従って、今回、構築したデータベースも時代の変化に合わせて、修正、追加だけでなく、新しいデータベースを構築する努力も必要になる。今後は、日々の業務を支援するデータベースをどう維持していくか、システム化することが重要である。

謝 辞

この調査研究は、平成 29 ～ 31 年度の富山県重点事業「病原体の遺伝子情報基盤整備による感染症対策強化事業」により遂行された。

参考文献

1. DDBJ (DNA Data Bank of Japan) , <https://www.ddbj.nig.ac.jp/services.html> (2020 年 7 月 1 日アクセス可能)
2. NCBI (National Center for Biotechnology Information) , <https://www.ncbi.nlm.nih.gov/nucleotide/> (2020 年 7 月 1 日アクセス可能)
3. Aziz RK, Bartels D, Best AA, et al. (2008) . BMC Genomics, 9:75
4. 秀丸エディタ , <https://hide.maruo.co.jp/software/hidemaru.html> (2020 年 7 月 1 日アクセス可能)
5. Carver T, Harris SR, Matthew Berriman M, et al. (2012) . Bioinformatics, 28 (4) , 464-469
6. Sequence Assistant Ver2.7, <http://www2s.biglobe.ne.jp/~haruta/> (2020 年 7 月 1 日アクセス可能)
7. Ruby script, <https://www.biostars.org/p/3722/> (2020 年 7 月 1 日アクセス可能)
8. FastQC, <http://www.bioinformatics.babraham.ac.uk/projects/fastqc/> (2020 年 7 月 1 日アクセス可能)
9. Bolger AM, Lohse M, Usadel B. (2014) . Bioinformatics, 30, 2114-2120
10. Marçais G, Delcher AL, Phillippy AM, et al. (2018) . PLoS Comput. Biol., 14, e1005944
11. Kozlov A, Darriba D, Flouri T, et al. (2019) . Bioinformatics, 35, 4453-4455
12. FigTree, <http://tree.bio.ed.ac.uk/software/figtree/> (2020 年 7 月 1 日アクセス可能)
13. NCBI Beta-Lactamase Data Resources, <https://www.ncbi.nlm.nih.gov/pathogens/beta-lactamase-data-resources/> (2020 年 7 月 1 日アクセス可能)
14. 地方衛生研究所全国協議会・国立感染症研究所. (2016). 病原体検出マニュアル 薬剤耐性菌, 2016 年 12 月版
15. Watahiki M, Kawahara R, Suzuki M, et al. (2020) . Jpn J Infect Dis., 73, 166-172
16. Tenover FC, Arbeit RD, Goering RV, et al. (1995) . J Clin Microbiol., 33, 2233-2239
17. Ooka T, Terajima J, Kusumoto M, et al. (2009) . J Clin Microbiol., 47, 2888-2894
18. Izumiya H, Pei Y, Terajima J, et al. (2010) . Microbiol.Immunol., 54, 569-577
19. Maiden MC, Bygraves JA, Feil E, et al. (1998) . Proc Natl Acad Sci USA, 95, 3140-3145
20. Octavia S, Lan R. J Clin Microbiol., (2007) . 45, 3795-3801
21. Walker TM, Clp CL, Harrell RH, et al. (2013) . Lancet Infect Dis., 13, 137-146
22. Maiden MC, Van Rensburg MJ, Bray JE, et al. (2013) . Nat Rev Microbiol., 11, 728-736