

The Genomes OnLine Database (GOLD) v.5: a metadata management system based on a four level (meta)genome project classifications

AUTHORS: T. B. K. Reddy^{1*}, Alex D. Thomas¹, Dimitri Stamatis¹, Jon Bertsch¹, Michelle Isbandi¹, Jakob Jansson¹, Jyothi Mallajosyula¹, Ioanna Pagani¹, Elizabeth A. Lobos¹ and Nikos C. Kyrpides^{1, 2*}

1. DOE Joint Genome Institute, 2800 Mitchell Drive, Walnut Creek, CA 94598, USA
2. Department of Biological Sciences, Faculty of Science, King Abdulaziz University, Jeddah, Saudi Arabia
- 3.

* To whom correspondence may be addressed. tbreddy@lbl.gov, and/or nckyrpides@lbl.gov

October 29, 2014

ACKNOWLEDGMENTS:

We are grateful to all the colleagues who kindly provide information for the more accurate monitoring of the genome projects and their associated metadata. We thank the members of the microbial genomics and metagenomics programs at the JGI for support, useful discussions and exchange of ideas. We would also like to thank Lynn Schriml and Lynette Hirschman for feedback and useful discussion on the EnvO ontology. This work was performed under the auspices of the US Department of Energy Office of Science, Biological and Environmental Research Program, and by the University of California, Lawrence Berkeley National Laboratory under contract no. DE-AC02-05CH11231.

DISCLAIMER:

This document was prepared as an account of work sponsored by the United States Government. While this document is believed to contain correct information, neither the United States Government nor any agency thereof, nor The Regents of the University of California, nor any of their employees, makes any warranty, express or implied, or assumes any legal responsibility for the accuracy, completeness, or usefulness of any information, apparatus, product, or process disclosed, or represents that its use would not infringe privately owned rights. Reference herein to any specific commercial product, process, or service by its trade name, trademark, manufacturer, or otherwise, does not necessarily constitute or imply its endorsement, recommendation, or favoring by the United States Government or any agency thereof, or The Regents of the University of California. The views and opinions of authors expressed herein do not necessarily state or reflect those of the United States Government or any agency thereof or The Regents of the University of California.

**Some supplementary files may need to be viewed online via your
Referee Centre at <http://mc.manuscriptcentral.com/nar>.**

The Genomes OnLine Database (GOLD) v.5: a metadata management system based on a four level (meta)genome project classification

T. B. K. Reddy^{1*}, Alex D. Thomas¹, Dimitri Stamatis¹, Jon Bertsch¹, Michelle Isbandi¹, Jakob Jansson¹, Jyothi Mallajosyula¹, Ioanna Pagani¹, Elizabeth A. Lobos¹ and Nikos C. Kyrpides^{1,2*}

¹ Prokaryotic Super Program, DOE Joint Genome Institute, Walnut Creek, CA 94598 USA

² Department of Biological Sciences, Faculty of Science, King Abdulaziz University, Jeddah, Saudi Arabia

* To whom correspondence should be addressed. Tel: 1-925-927-2580; Fax: 1-925-296-5666; Email: tbreddy@lbl.gov, nckyrpides@lbl.gov

ABSTRACT

The Genomes OnLine Database (GOLD, <http://www.genomesonline.org>) is a comprehensive online resource to catalogue and monitor genetic studies worldwide. GOLD provides up-to-date status on complete and ongoing sequencing projects along with a broad array of curated metadata. Here we report version 5 (v.5) of the database. The newly designed database schema and web user interface supports several new features including the implementation of a four level (meta)genome project classification system and a simplified intuitive web interface to access reports and launch search tools. The database currently hosts information for about 19,200 studies, 56,000 Biosamples, 56,000 sequencing projects, and 39,400 analysis projects. More than just a catalogue of worldwide genome projects, GOLD is a manually curated, quality controlled metadata warehouse. The problems encountered in integrating disparate and varying quality data into GOLD are briefly highlighted. GOLD fully supports and follows the Genomic Standards Consortium (GSC) Minimum Information standards.

INTRODUCTION

The Genomes OnLine Database (GOLD) is a data management system for cataloguing and continuous monitoring of sequencing projects worldwide. GOLD collects, curates, and disseminates metadata associated with those projects. GOLD is currently in its fifth version (1–6). With rapidly decreasing costs for sequencing, the number of sequencing projects and the amount of sequence data generated is increasing at an exponential rate. As these data are submitted to various public resources like GenBank (7) and EMBL (8) or analysis platforms like Integrated Microbial Genomes (IMG) (9) and MG-RAST (10), it becomes increasingly important to document the associated metadata in order to facilitate comparative analysis and hypothesis generation. The Genomics Standards Consortium (GSC) mandates the Minimum Information about any (x) Sequence (MlxS) specifications to be used when making sequence data available in public repositories (11, 12). GOLD is fully compliant with the GSC’s MlxS standards in capturing metadata and provides a platform to query projects based on various metadata features.

GOLD supports the IMG family of data management systems (9, 13–15) as a gatekeeper of projects and metadata, and requires that projects are annotated with at least minimal metadata. In fact an entry in GOLD and compliance with required metadata is a prerequisite to submit a project to the IMG systems for annotation. The main steps in the process include project registration in GOLD, project submission to IMG for annotation, and finally publication of results in the GSC's journal, Standards in Genomic Sciences (SIGS) (<http://www.standardsingenomics.com/>), or other journals of your choice. Since GOLD complies with MlxS, all available required metadata is already in place to publish in SIGS.

In the past, when sequencing was still expensive and only a limited number of high-interest organism genomes were sequenced, maintaining the associated information in a catalogue format was sufficient. With lower sequencing costs, many more genomes are now being sequenced as part of a single study. Initiatives such as the Human Microbiome Project (HMP) (16) and Genomic Encyclopedia of Bacteria and Archaea (GEBA) (17, 18) are a couple of examples where several thousands of genomes were sequenced as part of a single initiative. The emergence of high-throughput sequencing technologies and the development of analysis tools for studying metagenomes has facilitated the rapid growth in metagenome studies as well. It is also becoming more common to use multiple sequencing approaches on the same sample(s), for example the Functional Encyclopedia of Bacteria and Archaea (FEBA) (19). In such cases it is important to collect common metadata pertaining to these samples and organize all of the samples under one or more relevant studies.

The increasing variety of sequencing and analysis projects need to be linked and tracked in a seamlessly integrated system. One of the major limitations of the previous versions of the database has been the assumption of a one-to-one relationship between related components. For example, the previous versions could not correlate multiple sequencing projects to a single sample. In the event an isolate genome and metagenome were derived from a single sample a separate record for each sequence would need to be created. Similarly the previous versions could not capture the multiple sequencing projects of a combined assembly nor was it possible to connect multiple analyses to a single sequence project. Another limitation was that all genome projects were designated as isolates, an incorrect assignment for a genome assembled from a metagenome. These issues necessitated a new mechanism to organize various components of sequencing studies.

NEW TO THIS RELEASE

Version 5 of the database is founded on a fundamentally redesigned schema to accommodate a four level project classification system (Fig. 1). The new classification system is comprised of Studies, Biosamples, Sequencing Projects (SPs), and Analysis Projects (APs). Studies constitute the highest level of classification in the system, containing Biosamples, SPs, and APs that are

1
2
3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60

part of a single initiative. GOLD’s Biosamples represent the physical isolate or environmental material from which genetic material is extracted for sequencing. GOLD’s Biosamples have no relation to NCBI BioSamples. GOLD’s SPs represent sequencing protocols such as whole genome sequencing, transcriptomes, metagenomes, metatranscriptomes, methylation sequencing etc. applied to Biosamples. APs are the analytical processes applied to the SPs. Multiple different assemblies or annotations of the same SPs, would result to multiple different APs with varying metadata that need to be captured. These four components are described in more detail below.

In addition to the four levels described above, one more entity has been introduced in the new schema to provide metadata information for the individual organisms. In the previous versions of the database, each sequencing project of an isolate organism, included both the metadata for the sequencing information and the organism in a single record. Increasingly, the genome of a single organism is being sequenced more than once, by different groups, making it inefficient to associate the same organism metadata individually with every different project. GOLD v.5 defines and curates the organism records with core taxonomy, environmental, and other metadata independently of their associated SPs. As a result, this entity can be used by all SPs without the need for curating and propagating redundant metadata. By doing so, v.5 now enables the identification of all the organisms with different but synonymous names.

Historically, the focus of the database was to provide a comprehensive coverage to all prokaryotic genomes and metagenomes. We are in the process of systematically integrating eukaryotic SPs into GOLD. Projects are introduced in the database from three main streams: (1) projects sequenced at the JGI are automatically added following a number of Quality Control (QC) checks; (2) projects submitted to the database from individual researchers around the world; and (3) projects available at the NCBI’s BioProject portal.

The previous versions of the database provided read-only project reporting system. This served user needs for accessing project information and searching for projects based on specific metadata. However, the user interface for project creation and curation was provided through a separate system called IMG-GOLD. The new version has enabled the seamless integration of these two formerly separated functions into a single resource.

Isolate genomes via their associated Biosamples are now classified using the same five-tier hierarchical classification system previously developed and implemented for metagenomes (20). Over 10,000 public isolate genomes have been classified accordingly. Over 9,000 isolate genomes have also been curated to add strain habitat classifications. This field refers to the specific habitat of the strain according to the strain isolation information, as opposed to the previous general habitat in the database which corresponds to the species. The controlled vocabulary of the strain habitat has been mapped to the hierarchical ecosystem classification. For example, there are 161 genomes for organisms with the classification path Host-associated

(ecosystem) -> mammals (ecosystem category)-> digestive system (ecosystem type) -> foregut (ecosystem subtype)-> rumen (specific ecosystem). The strain habitats within this group include 'sheep rumen', 'cattle rumen' and 'goat rumen' with 37, 53, and 1 genome respectively. Thus, there is manual curation of organism Biosamples with specific habitat terms.

A newly designed web interface provides access to data through various pre-selected reports, project distribution graphs, statistics and an intuitive search interface that allows a user to search based on an array of metadata fields. The new implementation also provides access for public users to search for APs submitted to the IMG systems.

GOLD DATABASE ORGANIZATION AND DATA OVERVIEW

The Four-level classification system:

The current release organizes genome, metagenome, and other sequencing projects into a system of four levels which are described below.

GOLD Study:

A Study represents the highest-level organization. Studies include one or more Biosamples and their associated SPs and Aps that have been grouped to investigate a related research topic of interest. For example, the HMP (16), GEBA (17, 18) and KMG (21) studies represent typical cases where researchers set out to explore a specific topic by sequencing thousands of samples. Studies like GEBA-MDM (22) and FEBA (19) applied several different sequencing strategies (e.g., isolate genomes, single-cell genomes, metagenomes and transcriptomes etc.) as part of a single study. Studies may be composed of one to hundreds of Biosamples from a wide range of ecological settings (Fig. 2). Each Biosample may also yield several different SPs, each of which may yield multiple APs (Fig. 3, Table 1). Study IDs are referred to as "Gs" IDs in the new system. A GOLD Study is analogous to the NCBI's umbrella BioProject, and may contain one or more NCBI BioSamples.

GOLD Biosample:

Biosamples provides a description of the individual environmental sample, from which the organism, or genetic material (DNA or RNA) was isolated for downstream SPs. There are two types of Biosamples, organisms and biomes (environmental samples). Historically, samples were either isolated organisms for whole genome sequencing or environmental samples for metagenomics. However, it is becoming increasingly common to apply multiple sequencing techniques to a single sample and thus initiating several different SPs from the same starting material. For example, from a single biosample, whole DNA can be extracted for a metagenome and a metatranscriptome SP, as well as cells for single-cell genome project (Fig. 2) (19). The

need to manage and organize this type of complexity has led to the creation of GOLD Biosamples, which are quite distinct from NCBI's Biosamples. While GOLD Biosamples are organized above the sequencing projects in order to provide linkage of multiple sequencing projects originating from the same physical sample, NCBI's Biosamples are associated with individual sequencing projects, providing metadata only for that sequencing project. NCBI's Biosamples are also used *in lieu* of BioProjects to represent individual sequencing projects as in the case of multi-isolate projects. GOLD Biosample IDs are represented as “Gb” IDs.

GOLD Sequencing Project (SP):

A number of technological advances have enabled an increasing diversity of SP types (Fig. 3, Table 2). SPs represent individual sequencing deliverables such as metagenomes, metatranscriptomes, 16S sequences, single-cell genome sequences, isolate transcriptomes, or isolate whole genome sequences. As mentioned above, material from one Biosample can be the basis for more than one SP. GOLD SP's are often connected to a single NCBI BioProject, which could lead to the misconception that there is a one to one analogy between them. NCBI's BioProjects represent a mixture of project types that include the umbrella or multi-isolate types that are more analogous to the GOLD's Studies. This lack of standardization in NCBI BioProjects is one of the data management challenges the new GOLD classifications aims to address. GOLD Sequencing Project IDs are represented as “Gp” IDs. Each sequencing project can contain one or more APs.

GOLD Analysis Project (AP):

APs represent individual data processing methodologies or approaches that are undertaken for a given SP. As the diversity of data processing and analysis (eg assembly, structural, and functional annotation) methods has increased, so has the diversity of APs (Fig. 3). More specifically, the data generated from a single SP may be processed through multiple different approaches, as researchers exploring various different assembly methods or the same assembly with different annotation parameters. As shown in Figure 1, a researcher may also generate a combined assembly from multiple SPs and submit the data for annotation as one AP. This is more common in the case of single-cell genome projects where sparse sequence data from two related single cells can result in a better assembly and thereby more coverage of the genome of the organism being studied. One of the major limitation of the previous systems was the inability to represent these complex APs with their parent SPs. The current release fills this unmet need in representing different APs. AP IDs are represented as “Ga” IDs, and there are currently six different types:

(a) **Default AP:** this represents the standard assembly and annotation process applied for any sequencing project.

(b) **Default-screened AP and default - unscreened AP**: these are applicable only for single cell genome projects where contamination is a major issue due to extraneous DNA or due to errors during cell sorting/isolation events. Accordingly, there is a need to distinguish between APs that have gone through a decontamination round (screened) versus those that have not (unscreened).

(c) **Combined assembly AP**: these APs use data from multiple SPs which are combined into a single assembly, which is then submitted for annotation. For example whole genome shotgun sequencing may be applied to a set of single cell genomes from the same Biosample and the data from each single cell genome can be used to generate a combined assembly for a better genome reconstruction. Alternatively, metagenomic sequences from multiple different Biosamples may be combined into a single assembly. Tracking these many relationships between Biosamples, SPs, and APs within a Study is a key feature of new GOLD.

(d) **Genome from metagenome AP**: these APs represent individual genomes extracted from metagenomics data. Advances in metagenomic assembly and binning (<http://ggkbase.berkeley.edu/>, 23) have enabled the reconstruction of partial or entire genomes directly from metagenomic sequencing project.

(e) **Reassembly AP**: represent the APs created when an already processed genome is subjected to different assembly methods to generate a new assembly.

(f) **Reannotation AP**: represent the AP created for annotating a genomes that has been annotated before.

(g) **Metatranscriptome Mapping AP**: these APs represent the mapping of the metatranscriptomic data on the metagenomic sequences in order to connect functional processes to genes.

GOLD BY NUMBERS

Studies:

As of September 2014, there are 19,242 Studies in GOLD. These include 472 metagenomic studies (i.e. have at least one metagenome sequencing project) and 18,770 non-metagenomic studies. Studies have been generally growing in size and complexity and are increasingly composed of Biosamples from more diverse environments (Fig. 2). There is also an increasing number of sequencing strategies applied to each Biosamples (Fig. 2, Table 1) as well a growing number of APs used within a study (Fig. 3).

Biosamples:

There are currently 56,403 Biosamples in the database which are classified as host-associated (9,922 samples), engineered (1,270 samples), environmental (3,637 samples), and unclassified (36,435 samples). Organism Biosamples represent more than 150 GOLD phylogenetic

classifications. Biome Biosamples represent more than 200 unique GOLD ecosystem classifications.

Sequencing Projects:

There are currently 56,458 Sequencing Projects reported in the database. These include 47,932 whole genome sequencing (WGS) projects distributed across 36,824 bacteria, 5,822 eukaryal, and 851 archaeal projects. There are also 4,351 metagenomic SPs, distributed across 1,567 host-associated, 239 engineered, and 2,545 environmental projects. In addition to the genomic and metagenomic SP, the database provides information on 1,200 transcriptomic and 797 metatranscriptomic SPs. While there are only 50 targeted gene survey SPs, all of these are part of Studies that include metagenomic data and most include metatranscriptomic data (Table 1). The database also provides information on 13 transposon mutagenesis SPs. As this technique is becoming more high-throughput more projects of this type can be expected (19). A similar growth is expected for the methylation SPs, only 15 of which are currently available in the database.

Analysis Projects:

38,573 APs are currently reported of which 36,755 are default APs. For single-cell SPs there are 856 default - screened and 1082 default-unscreened APs. There are also 107 transcriptome mapping and 80 metatranscriptome mapping APs. Finally, 30 combined assembly APs from 310 SPs in 11 Studies are available in GOLD. All of the Sequencing Projects used for combined assembly were also used for 'default' APs.

ACCESSING GOLD

GOLD provides free access to all publicly available data, project status reports and other statistical information. Data can be accessed by various pre-computed reports or by querying the database using search functions. Menu tabs to allow users to choose Search, Distribution Graphs, Biogeographical Metadata and Statistics options to access data are also available from the front page. A list of all the public projects in the database is also available for download.

Distribution Graphs:

Automatically generated pie charts that describe the different types of projects in the database are now available. These include data organized by SP type, sequencing status, phylogenetic table, phylogenetic tree and Biosample classification in separate tabs.

Biogeographical Metadata:.

The geographic distribution of Biosamples can be visualized via the Google Map and Google Earth options. These can also be used to select Biosamples based on their geographic location.

The Google Map feature aggregates Biosamples by geographic location into circles noting the number of Biosamples in a group when viewing larger spatial extents. These groupings are ungrouped as the map view is focused using the zoom feature. The map view can be focused on the location of a biosample when it is selected from a list next to the map. The Google Earth feature provides a similar tool but with a 3-dimensional global perspective.

Statistics:

The GOLD statistics page provides several precomputed user friendly, easy to interpret graphs, bar charts and pie charts about various sequencing projects. Refer to supplementary material for more details about various precomputed charts.

SEARCHING THE GOLD DATABASE

The Search function can be used to query the database based on various search criteria that encompass all four levels of the project classification system or based on various metadata features. A drop-down menu allows the choice of three search options, Quick Search, Advanced Search, and Metadata Search.

Quick Search:

Quick Search allows a user to search through the most frequently used fields/identifiers across the four levels in the database (Studies, Biosamples, Sequencing Projects and APs).

Advanced Search:

Advanced Search provides options to query metadata fields in each level of the new classification system. Results are provided as a list according to the search criteria, with fields used displayed in separate columns. Search result can be redefined by removing any search term by clicking 'remove' next to the search term in the column header. Search results may also be refined directly in the results table by modifying the search term to any field by clicking the "+" under the column header. There is also a 'Select Fields' button on the left, which allows the user to add additional fields.

Metadata Search:

Metadata search is designed to query the database using various metadata identifiers. These include the classification by the domains of the project organism, Archaea, Bacteria, Eukarya or all. The various search tabs contain graphical and tabular representation of the numbers of projects or organisms. This approach serves to obtain an overall picture of projects and samples according to chosen criteria, and produces a sortable table and also plots these lists in a pie-chart for easy reference.

CREATING AND EDITING PROJECTS IN GOLD

Registered users can submit new projects or edit their existing entries.

Editing: Existing projects can be updated using a new inline-editing user interface. For editing existing entries a user needs to login and select the entry of their interest. When clicking on a field, an edit box is launched with existing values in it. One can update the value and save. The inline-edit feature seamlessly integrates the edit functionality with user interface without the need for launching a separate edit form.

Creating New Projects: Registered users can create new SPs using the new project entry interface. Creating a new SP also requires defining all related database entities i.e. Study, Biosample, and Organism when applicable for isolate genome projects. As shown in supplementary material the new project entry landing page provides the following options (a) create a new SP (b) create new AP (c) review your Studies, Biosamples, and Sequencing Projects.

The new SP creation interface will walk a user through a series of steps to define new projects or select existing projects. For example, launching ‘Create a new Sequencing Project’ will first ask if this is metagenome (biome) or isolate (organism) project. This information is used to launch appropriate forms and guide users through the process. Next, a user will be asked to enter a Study for the SP. If this is a returning user adding additional SPs to an existing Study, the user will be able to choose the existing Study. Otherwise the user will be asked to define a new one. Once the Study is created a Biosample must be defined. Again the user may define a new Biosample or select an existing Biosample. If the SP is for an isolate organism, the user must select an existing organism from the database or define a new organism. After the Study, Biosample and/or organism are created, the user will be able to define a new SP. All the required fields are marked with an asterisk and tool tips are provided with appropriate examples to guide a user in defining new projects. Help pages are available to provide explanation on specific database terminology. If a SP is defined a user can select “Create a new Analysis Project for submission to IMG” to define an AP. A single SP can have multiple APs to represent different assemblies and/or gene calling methodologies applied. Study, Biosample and Organism entries created but not yet associated with a sequencing project are saved as drafts. Users can access these from the “My Data” table as well as select these from the pulldown list as part of new SP creation interface.

DATA IMPORT AND CURATION CHALLENGES

GOLD continuously monitors sequencing projects around the world both through direct submissions from users and through data imports from major public resources, such as NCBI (7). A series of cross checks have been implemented to ensure high data quality, manually verify data

conflicts, and curate metadata during and after import into the database. Due to the nature of data organization and data quality enforcement standards at different public resources it is challenging and curation intensive to keep the import processes working. For a list of examples see the supplementary material.

The aim of listing these issues is two fold (i) to express the difficulties that any integrated public database resource like GOLD is facing in representing disparate information and (ii) to highlight the need for more manual data curation and quality control checks at major public resources like NCBI. If the data are corrected at the source it saves time and effort for several groups around the world. For example correctly representing the sequencing center names and geographic coordinates at the source would eliminate the need for all other databases who import data from NCBI to come up with their own procedures for finding and resolving these issues. NCBI systems serve as a large democratizing force providing unrestricted access for users around the world to submit their data and freely share with the rest of the world. With such a broad mandate and unhindered access, it is difficult to enforce strict standards, but at least some the above listed issues can be mitigated with more manual curation and quality control processes in place. These challenges are not unique to this database, but to all who rely on public database resources. Thus there is a strong case for the stakeholders and funding agencies to support data curation efforts at public resources (23).

FUTURE DIRECTIONS:

Future developments will focus on data integration, expanding metadata fields and providing sophisticated search options across the metadata fields at different classification levels.

Data Integration: we will continue importing public metagenome sequencing projects from NCBI and EBI. We will expand our semi-automatic NCBI isolate genome import process to include multi-isolate NCBI BioProjects, where more than one isolate genome is listed under a single NCBI BioProject with different NCBI BioSamples as opposed to represented by individual NCBI BioProjects.

Expanding metadata fields: The growing complexity of the SPs and the diversity of the GOLD Biosamples collected from specific locations and conditions necessitate GOLD to constantly expand metadata fields. We plan to incorporate all of the MlxS environmental packages and include metadata fields that are not currently available in the database.

Metadata Miner: The advanced search feature in the current release provides an option to search among a multitude of metadata fields within each of the four project classification levels. For data mining and hypothesis generation often it is important to search across different levels using

1
2
3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60

different metadata fields at the same time. For example the search for “aerobic bacterial whole genome sequencing projects that have a project relevance of medical, with human as a host and project status of complete” in the current implementation would need to be executed in multiple steps at different GOLD classification levels. We plan to implement an integrated Metadata Miner that would facilitate complex searches across all four levels of GOLD. Such an advanced metadata mining tool will make it easy for users to execute searches similar to the above example.

CONCLUSION

The steady increase in the number of sequencing studies carried out around the world coupled with the complexity of the samples, diversity of sequencing strategies, and expanding analysis methods necessitates an integrated metadata warehouse like GOLD. As outlined above both through our current release and proposed feature enhancements like Metadata Miner, GOLD is uniquely positioned to organize sequence metadata and provide unhindered access both for hypothesis generation and testing. GOLD’s rich metadata coupled with seamless integration with the IMG analysis systems provides users with the ability to look at their data and analyze results from a whole different perspective with associated metadata. This helps in understanding the observations as well as asking questions to find answers hitherto impossible without curated metadata. Towards this goal GOLD will continue expanding in terms of metadata fields well as the numbers of projects integrated from various sources around the world.

ACKNOWLEDGEMENTS

We are grateful to all the colleagues who kindly provide information for the more accurate monitoring of the genome projects and their associated metadata. We thank the members of the microbial genomics and metagenomics programs at the JGI for support, useful discussions and exchange of ideas. We would also like to thank Lynn Schriml and Lynette Hirschman for feedback and useful discussion on the EnvO ontology.

FUNDING

This work was performed under the auspices of the US Department of Energy Office of Science, Biological and Environmental Research Program, and by the University of California, Lawrence Berkeley National Laboratory under contract no. DE-AC02-05CH11231.

REFERENCES

1. Kyrpides, N.C. (1999) Genomes OnLine Database (GOLD 1.0): a monitor of complete and ongoing genome projects world-wide. *Bioinformatics*, **15**, 773–774.

2. Bernal, a, Ear, U. and Kyrpides, N. (2001) Genomes OnLine Database (GOLD): a monitor of genome projects world-wide. *Nucleic Acids Res.*, **29**, 126–7.
3. Liolios, K., Tavernarakis, N., Hugenholtz, P. and Kyrpides, N.C. (2006) The Genomes On Line Database (GOLD) v.2: a monitor of genome projects worldwide. *Nucleic Acids Res.*, **34**, D332–4.
4. Liolios, K., Mavromatis, K., Tavernarakis, N. and Kyrpides, N.C. (2008) The Genomes On Line Database (GOLD) in 2007: status of genomic and metagenomic projects and their associated metadata. *Nucleic Acids Res.*, **36**, D475–9.
5. Liolios, K., Chen, I.-M. a, Mavromatis, K., Tavernarakis, N., Hugenholtz, P., Markowitz, V.M. and Kyrpides, N.C. (2010) The Genomes On Line Database (GOLD) in 2009: status of genomic and metagenomic projects and their associated metadata. *Nucleic Acids Res.*, **38**, D346–54.
6. Pagani, I., Liolios, K., Jansson, J., Chen, I.-M. a, Smirnova, T., Nosrat, B., Markowitz, V.M. and Kyrpides, N.C. (2012) The Genomes OnLine Database (GOLD) v.4: status of genomic and metagenomic projects and their associated metadata. *Nucleic Acids Res.*, **40**, D571–9.
7. Benson, D. a, Cavanaugh, M., Clark, K., Karsch-Mizrachi, I., Lipman, D.J., Ostell, J. and Sayers, E.W. (2013) GenBank. *Nucleic Acids Res.*, **41**, D36–42.
8. Leinonen, R., Akhtar, R., Birney, E., Bower, L., Cerdeno-Tárraga, A., Cheng, Y., Cleland, I., Faruque, N., Goodgame, N., Gibson, R., et al. (2011) The European Nucleotide Archive. *Nucleic Acids Res.*, **39**, D28–31.
9. Markowitz, V.M., Chen, I.-M. a, Palaniappan, K., Chu, K., Szeto, E., Pillay, M., Ratner, A., Huang, J., Woyke, T., Huntemann, M., et al. (2014) IMG 4 version of the integrated microbial genomes comparative analysis system. *Nucleic Acids Res.*, **42**, D560–7.
10. Meyer, F., Paarmann, D., D'Souza, M., Olson, R., Glass, E.M., Kubal, M., Paczian, T., Rodriguez, a, Stevens, R., Wilke, a, et al. (2008) The metagenomics RAST server - a public resource for the automatic phylogenetic and functional analysis of metagenomes. *BMC Bioinformatics*, **9**, 386.
11. Field, D., Amaral-Zettler, L., Cochrane, G., Cole, J.R., Dawyndt, P., Garrity, G.M., Gilbert, J., Glöckner, F.O., Hirschman, L., Karsch-Mizrachi, I., et al. (2011) The Genomic Standards Consortium. *PLoS Biol.*, **9**, e1001088.
12. Yilmaz, P., Kottmann, R., Field, D., Knight, R., Cole, J.R., Amaral-Zettler, L., Gilbert, J. a, Karsch-Mizrachi, I., Johnston, A., Cochrane, G., et al. (2011) Minimum information about a marker gene sequence (MIMARKS) and minimum information about any (x) sequence (MInXS) specifications. *Nat. Biotechnol.*, **29**, 415–20.
13. Markowitz, V.M., Chen, I.-M. a, Chu, K., Szeto, E., Palaniappan, K., Pillay, M., Ratner, A., Huang, J., Pagani, I., Tringe, S., et al. (2014) IMG/M 4 version of the integrated metagenome comparative analysis system. *Nucleic Acids Res.*, **42**, D568–73.
14. Markowitz, V.M., Mavromatis, K., Ivanova, N.N., Chen, I.-M. a, Chu, K. and Kyrpides, N.C. (2009) IMG ER: a system for microbial genome annotation expert review and curation. *Bioinformatics*, **25**, 2271–8.

1
2
3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60

15. Markowitz, V.M., Chen, I.-M. a, Chu, K., Szeto, E., Palaniappan, K., Jacob, B., Ratner, A., Liolios, K., Pagani, I., Huntemann, M., et al. (2012) IMG/M-HMP: a metagenome comparative analysis system for the Human Microbiome Project. *PLoS One*, **7**, e40151.

16. Nelson, K.E., Weinstock, G.M., Highlander, S.K., Worley, K.C., Creasy, H.H., Wortman, J.R., Rusch, D.B., Mitreva, M., Sodergren, E., Chinwalla, A.T., et al. (2010) A catalog of reference genomes from the human microbiome. *Science*, **328**, 994–9.

17. Wu, D., Hugenholtz, P., Mavromatis, K., Pukall, R., Dalin, E., Ivanova, N.N., Kunin, V., Goodwin, L., Wu, M., Tindall, B.J., et al. (2009) A phylogeny-driven genomic encyclopaedia of Bacteria and Archaea. *Nature*, **462**, 1056–60.

18. Kyrpides, N.C., Hugenholtz, P., Eisen, J. a, Woyke, T., Göker, M., Parker, C.T., Amann, R., Beck, B.J., Chain, P.S.G., Chun, J., et al. (2014) Genomic Encyclopedia of Bacteria and Archaea: Sequencing a Myriad of Type Strains. *PLoS Biol.*, **12**, e1001920.

19. Blow, M.J., Deutschbauer, A.M., Hoover, C.A., Lamson, J., Pennacchio, L.A., Price, M.N., Waters, J., Wetmore, K.M., Bristow, J. and Arkin, A.P. (2013) Functional Encyclopedia of Bacteria and Archaea. Poster session presented at: Genomics of Energy & Environment User Meeting; Walnut Creek, CA.

20. Ivanova, N., Tringe, S.G., Liolios, K., Liu, W.-T., Morrison, N., Hugenholtz, P. and Kyrpides, N.C. (2010) A call for standardized classification of metagenome projects. *Environ. Microbiol.*, **12**, 1803–5.

21. Kyrpides, N.C., Woyke, T., Eisen, J. a, Garrity, G., Lilburn, T.G., Beck, B.J., Whitman, W.B., Hugenholtz, P. and Klenk, H.-P. (2014) Genomic Encyclopedia of Type Strains, Phase I: The one thousand microbial genomes (KMG-I) project. *Stand. Genomic Sci.*, **9**, 1278–84.

22. Rinke, C., Schwientek, P., Sczyrba, A., Ivanova, N.N., Anderson, I.J., Cheng, J.-F., Darling, A., Malfatti, S., Swan, B.K., Gies, E. a, et al. (2013) Insights into the phylogeny and coding potential of microbial darkmatter- Supplementary Information. *Nature*, **499**, 431–7.

23. Kyrpides, N.C. (2009) Fifteen years of microbial genomics: meeting the challenges and fulfilling the dream. *Nat. Biotechnol.*, **27**, 627–32.

TABLE AND FIGURE LEGENDS

Table 1. GOLD Sequencing strategy combinations used within a study.

Figure 1. Four level project classification system implemented in v.5 to describe studies, Biosamples, sequencing projects, and analysis projects. Studies group one or more related Biosamples. Biosamples describe an individual sample of genetic material. Sequencing projects are the sequencing deliverables from the Biosamples. Analysis projects are the data processing methods applied to sequencing projects. A) Biosamples may be merged prior to sequencing projects (ex. 16S amplicon data combined prior to sequencing). B) Sequencing Projects may be merged prior to analysis (ex. multiple single-cell genomes combined for assembly).

Figure 2. Study Biosamples, ecosystem categories, and sequencing strategies. Each point is a GOLD study. The size of the point represents the number of ecosystem categories within a study. The position on the y-axis notes the number of Biosamples within a study. The color of each point notes the number of unique sequencing strategies used within a study.

Figure 3. Sequencing and analysis projects per study over time. Color notes the types of sequencing strategies used within a study. Size indicates the number of analysis projects within a study.

1
2
3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60

Table 1. GOLD Sequencing strategy combinations used within a study.

Sequencing Strategy Combinations	No. Studies	No. Sequencing Projects
Whole Genome Sequencing	18211	46905
Metagenome	403	3315
Transcriptome, Whole Genome Sequencing	76	1989
Metagenome, Metatranscriptome	39	927
Metagenome, Metatranscriptome, Targeted Gene Survey	6	682
Transcriptome	402	596
Metagenome, Whole Genome Sequencing	1	217
Metatranscriptome	9	54
Metagenome, Targeted Gene Survey	4	53
smRNA, Transcription Start Site, Transcriptome, Transposon Mutagenesis Sequencing, Whole Genome Sequencing	1	34
Metatranscriptome, Targeted Gene Survey	1	21
Methylation	4	15
smRNA, Transcriptome	1	14
Plasmid	2	2
smRNA	1	1
Transposon Mutagenesis Sequencing	1	1

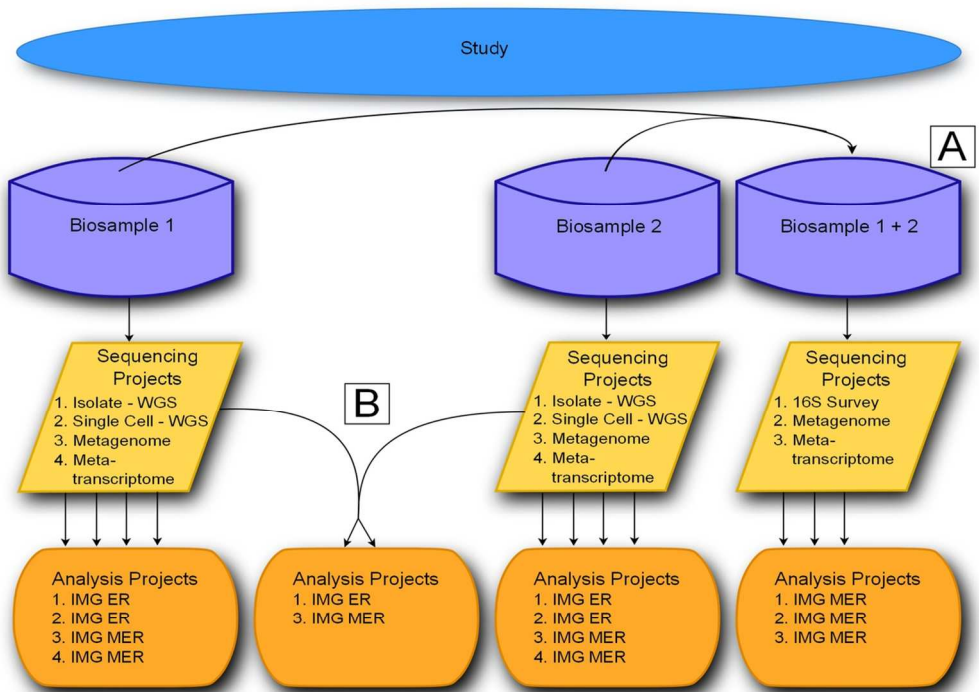


Figure 1
117x84mm (300 x 300 DPI)

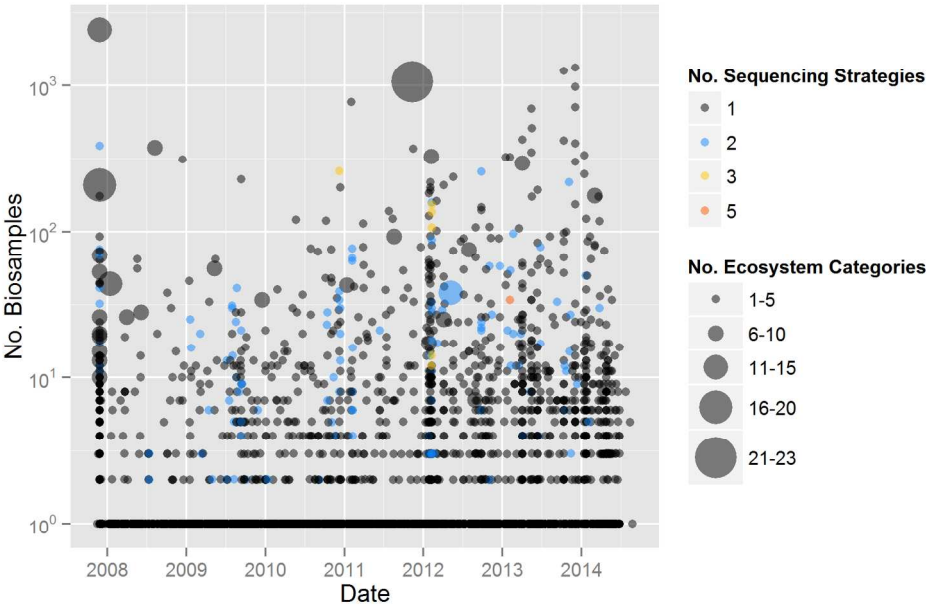


Figure 2

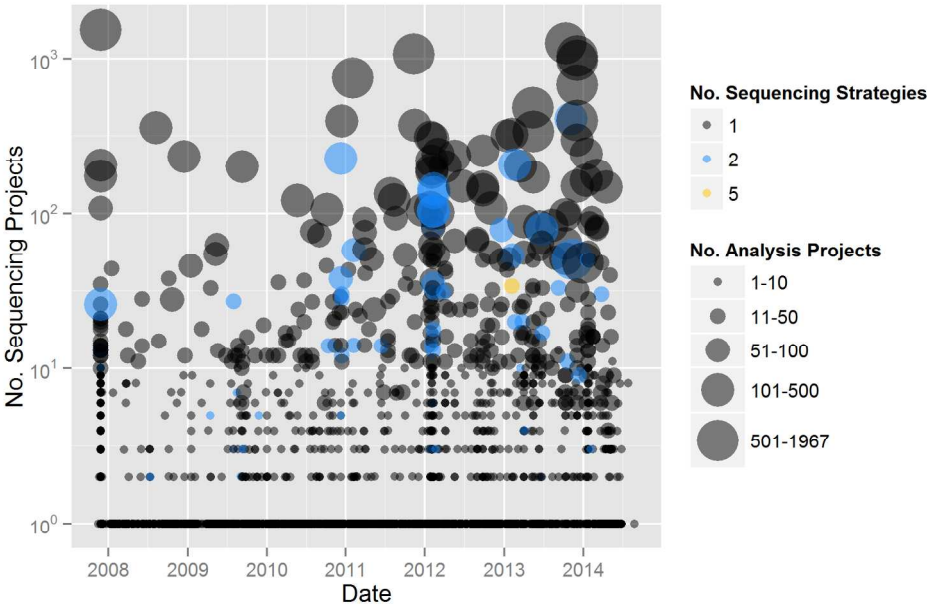


Figure 3

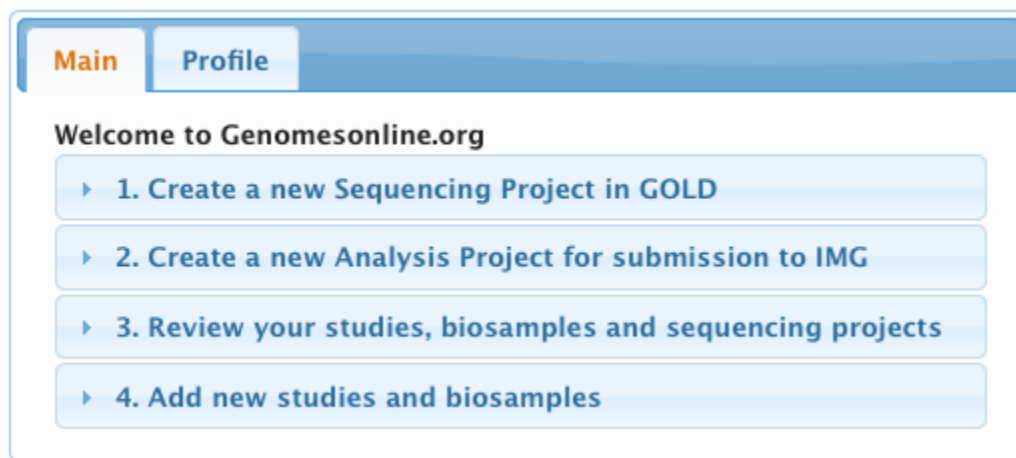
1
2
3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60

Supplementary Materials

S1. STATISTICS:

- i) Complete and Permanent Draft Genome Totals (by year and status): This bar chart lists Complete and Permanent Draft genome projects that were completed over the last 9 years.
- ii) Genome Totals (by year and Status): This line diagram tracks the total number of genome sequencing projects registered for each year since 1995. These are further divided into complete and incomplete projects.
- iii) Project Totals (by year and Phylogenetic Group): This graph represents Archaeal, Bacterial, Eukaryotic and Metagenomic projects registered for every year since 2007.
- iv) Phylogenetic distribution of Bacterial Genome Projects: This pie chart breaks down bacterial genome sequencing projects according to major bacterial phyla like Actinobacteria, Firmicutes and Proteobacteria. The number of genomes sequenced in each category are represented in this pie chart.
- v) Projects By Major Sequencing Centers: This pie chart displays the number of genomes sequenced by major sequencing centers worldwide. It may be noted that the top seven major sequencing centers represent 73% of the total projects sequenced.
- vi) Project Relevance of Bacterial Genome Projects: This pie chart shows the breakdown of the number of projects annotated with different relevance terms such as Medical, Environmental, Biotechnological etc.
- vii) Major Sequencing Centers for Archaeal and Bacterial Genomes: This pie chart shows the number of Archaeal and Bacterial genomes sequenced by major sequencing centers.

S2. CREATING AND EDITING PROJECTS IN GOLD



New project entry landing page for creating new projects in GOLD

S3. DATA IMPORT AND CURATION CHALLENGES

Here is a select list of issues we face while importing projects into GOLD.

- Frequent data format changes. This necessitates investigating the changes vis-a-vis their overall impact on data import processes in GOLD as well as updating the necessary scripts.
- Ambiguous representation of Genome Projects at BioProject site. Eg: NCBI BioProject ID: 189730. This is a whole genome sequencing project. But it is annotated as Material = Other and Capture = Other instead of correctly representing it as Material = Genome; Capture = Whole. So we can't import such projects automatically. We have to manually review and determine if they are sequencing projects or not.
- Multi-isolate projects. Instead of creating a separate BioProject entry for each isolate genome, NCBI often lists multiple isolate genomes under a single BioProject. Though they are labelled as Scope = "Multiisolate", it require special handling to represent them as in GOLD as individual genome projects. Eg: NCBI BioProject ID: 238952 consists of 6 different individual isolate genomes. In GOLD we represent these as 6 different sequencing projects. As such there is no error on the part of NCBI. This is a classic example of how different resources organize and represent data in their systems. Though such differences look subtle, it often require both engineering and curation efforts to make data imports work.
- We often encounter projects that were listed as multi-isolates, but they are not actually multi-isolate projects. We manually check all multi-isolate projects before importing them into GOLD. It is possible BioProject users specify these erroneously as multi-isolates during submission. Eg: PRJNA238854 appears to be reporting an isolate genome but it is scoped as multi-isolate.
- Cryptic or uninformative sequencing center names on BioProjects. Sequencing center name is a required field in GOLD. We use this info to track sequencing projects carried out at major sequencing centers and provide related statistics on GOLD.
Eg: UH, UB, TDC, 'wqedd wed we' etc. BioProject ID: 34783 lists 'wqedd wed we' as sequencing center name. Again these values are what users provided at the time of project submission to BioProject. We look at the GenBank record or associated publications to infer what UH, UB, TDC etc. mean and curate the same in GOLD.

- Inaccurate geolocation info due geolocation name and coordinates mismatch.

Eg: BioSamples under PRJNA252784 and PRJNA252785 list China, Liaoning as geographic location. Where as the coordinates map to Northern California, USA.

- Phage genome sequence included as part of host genome sequence. Eg: PRJNA247.

Chlamydomonas reinhardtii AR39 genome project also listing Chlamydia phage phiCPAR39 genome sequencing.

Duplicate BioProjects: PRJNA243100 is a duplicate entry for PRJNA243545. Both entries are for the same organism from same institute Shanghai JiaoTong University. In GOLD we choose to list PRJNA243545 which is associated with GenBank record.