

Table des matières

INTRODUCTION GENERALE	1
CHAPITRE I	
L'ETUDE PREALABLE.....	5
I. INTRODUCTION	6
II. SGBDs	6
III. OUTILS DE MANIPULATION DE DONNEES	7
<i>III.1 Manipulation des données.....</i>	<i>7</i>
<i>III.2 Data Profiling</i>	<i>13</i>
<i>III.3 Génération de données.....</i>	<i>15</i>
<i>III.4 Etude Comparative.....</i>	<i>20</i>
<i>III.5 Conclusion</i>	<i>22</i>
CHAPITRE II	
ANALYSE ET CONCEPTION	23
I. INTRODUCTION	24
II. PROCESSUS DE DEVELOPPEMENT	24
III. SPRINT.....	26
IV. ARCHITECTURE.....	27
V. CONCEPTION AVEC UML	28
<i>V.1 Diagramme de cas d'utilisation.....</i>	<i>29</i>
<i>V.2 Diagramme de séquence.....</i>	<i>30</i>
<i>V.3 Diagramme de classe.....</i>	<i>36</i>
VI. CONCEPTION DE L' AUTO-COMPLETION	37
VII. CONCLUSION	40
CHAPITRE III	
IMPLEMENTATION	41
I. INTRODUCTION	42
II. REALISATION.....	42
III. GESTION DES TESTS ET ANOMALIES.....	53
IV. OUTILS DE DEVELOPPEMENT.....	54
<i>IV.1 Gestion de projet « Trello ».....</i>	<i>54</i>
<i>IV.2 Modélisation « astah »</i>	<i>55</i>
<i>IV.3 Programmation « Embarcadero Delphi ».....</i>	<i>55</i>
V. PROGRAMMATION	56
VI. CONCLUSION.....	56
CONCLUSION GENERAL.....	57
REFERENCES BIBLIOGRAPHIQUES.....	59
LISTE DE FIGURES	61
LISTE DES TABLEAUX.....	62
LISTE DES ABREVIATIONS	63

INTRODUCTION GENERALE

Contexte

Notre travail s'intègre dans le contexte des Systèmes de Gestion de Base de Données (SGBD). Ces derniers représentent un ensemble de services (applications logicielles) permettant de gérer les bases de données, c'est à dire : l'accès simple, autoriser un accès aux informations à de multiples utilisateurs, manipuler les données présentées dans la base de données (insertion, suppression, modification) [1].

Les systèmes de gestion de base de données sont des logiciels universels, indépendants de l'usage qui est fait des bases de données. Ils sont utilisés pour plusieurs applications informatiques, comme les guichets automatiques bancaires, les logiciels de réservation, les bibliothèques numériques, les logiciels d'inventaire, les logiciels de gestion intégrés et sites web. Il existe de nombreux systèmes de gestion de base de données. En 2008, Oracle détenait près de la moitié du marché des SGBD avec MySQL et Oracle Database. Ensuite IBM est venu avec près de 20 %, laissant peu de place pour les autres acteurs. Les SGBD sont souvent utilisés par d'autres logiciels ainsi que les administrateurs ou les développeurs. Ils peuvent être sous forme de composants logiciels, de serveurs, de logiciels applicatifs ou d'environnement de programmation. Un système de gestion de base de données permet de faire le stockage et le partage des informations dans une base de données. En garantissant la qualité, la pérennité et la confidentialité des informations, en cachant bien sur la complexité des opérations. Les données sauvegardées dans des bases de données modélisent des objets du monde réel, ou des associations entre objets. Les objets sont en général représentés par des articles de fichiers, alors que les associations correspondent naturellement à des liens entre articles. Les données peuvent donc être vues comme un ensemble de fichiers par des pointeurs ; elles sont interrogées et mises à jour par des programmes d'application écrits par les utilisateurs ou par des programmes utilitaires fournis avec le SGBD. Un SGBD peut être défini comme un langage (logiciel) qui sert à dialoguer avec une BDD et qui permet à l'utilisateur de définir les données d'une base, de les consulter et de les mettre à jour [2].

Environnement de travail :

Ce travail a été réalisé dans le contexte d'un stage dans la société Soft Builder qui est une entreprise spécialisée dans les outils d'administration de SGBD. Elle a été fondée en 2015. Cette société développe un produit de modélisation de données appelé ER-One. Notre travail consiste à développer un outil qui le complète pour offrir un logiciel complet d'administration de SGBD.

Le stage s'est déroulé du 30 janvier au 29 juin 2017. Nous avons choisi cette entreprise, car nous avons constaté que le thème complète les connaissances que nous avons acquises durant nos études à l'université.

Problématique

La diversité et la disponibilité de plusieurs logiciels de manipulation, traitement et gestion de base données répond à de nombreux besoins et problèmes. Malgré le développement de ces derniers et leurs avantages, et vu que l'utilisateur est ambitieux et il cherche toujours à trouver les meilleures solutions, il est nécessaire de faire évoluer ces outils. Nous avons fait une étude sur plusieurs outils d'administration et de manipulation de données et nous avons remarqué que la plupart des outils sont dédiés à un seul SGBD, ils gèrent une seule connexion active à une BDD à la fois et ils n'offrent pas des rapports dynamiques (navigables). Nous avons constaté qu'il n'existe pas un outil qui propose toutes les fonctionnalités de gestion et d'administration et manipulation de données. Le développeur ou l'administrateur de BDD utilise plusieurs applications.

La société Soft-Builder veut concevoir et offrir une seule solution intégrée afin de répondre à ces problèmes. Cette société a déjà un outil appelé ER-One qui permet la création des modèles entité/relation à l'aide d'un éditeur visuel, la création automatique des bases de données sous Oracle, MySQL, SQL Server, etc., le reverse engineering, etc.

Contribution

Notre travail consiste à réaliser un module de manipulation de données multi SGBD, qui sera intégré à ER-One et qui offrira les fonctionnalités suivantes :

- navigations dans les données, leurs mise à jour, import/export (CSV, SQL et XML), et aide à la création de requêtes SQL.

- Connexion à plusieurs SGBD : Oracle, MySQL, PostgreSQL, SQL Server et Firebird.
- Génération des données avec différentes méthodes en définissant le nombre de ligne à générer et en donnant un type pour chaque champ selon le besoin de l'utilisateur (dictionnaire, séquence, aléatoire).
- Visualisation et export des rapports qui contiennent des statistiques concernant une table ou une base de données précise « Data Profiling ».
- Gestion des données (ajouter, supprimer, modifier).
- Trie, duplication et filtrage les données.
- Exécution des requêtes SQL.

Plan du mémoire

Notre mémoire est structuré comme suit :

Chapitre 1 : présente l'étude des SGBDs, des outils de manipulation de données, ainsi que la génération de données et le data Profiling. En fin, nous terminerons par une étude comparative.

Chapitre 2 : présente la planification et la démarche de notre travail qui donne un aperçu détaillé de la conception de notre application.

Chapitre 3 : présente l'implémentation et les tests qui donnent un aperçu détaillé de l'application que nous avons élaboré.

Nous terminerons ce document par une conclusion générale qui résume notre travail avec des perspectives pour les futurs travaux de ce projet.

CHAPITRE I

L'ETUDE PREALABLE

I. Introduction

Comme nous l'avons déjà signalé dans l'introduction générale, le but de notre projet est de réaliser un outil qui se connecte à plusieurs SGBD pour faciliter leur utilisation et leur gestion en un seul outil. Dans ce chapitre nous commençons par présenter les SGBD qui seront supportés par notre application. Ensuite nous allons présenter quelques outils similaires à notre système que nous avons scindé en trois parties : manipulation de données, data profiling et génération de données.

II. SGBDs

Puisque nous avons réalisé l'accès à 5 SGBD dans notre application, nous allons présenter :

MySQL :

Base de données libre très utilisée par les hébergeurs web pour ses bonnes performances et dont la version 5 comble les principaux manques (trigger, procédures stockées...). MySQL a la particularité de fournir plusieurs moteurs de bases de données (InnoDB, MyIsam, Berkeley DB...), qui ne fournissent pas les mêmes fonctionnalités ni les mêmes performances, afin de s'adapter au besoin. Par exemple, le moteur ARCHIVE est optimisé pour sauvegarder et fournir des informations, mais il n'est pas possible de mettre à jour une ligne dans une table [3].

Oracle :

Est un système de gestion de base de données relationnelle (SGBDR) qui depuis l'introduction du support du modèle objet dans sa version 8 peut être aussi qualifiée de système de gestion de base de données relationnel-objet (SGBDRO) [4].

PostgreSQL :

PostgreSQL est un système de SGBDRO (gestion de base de données relationnelle et objet). Il fonctionne selon une architecture client/serveur. Le coté serveur, PostgreSQL marche sur la machine en hébergeant la BDD qui est capable de traiter les requêtes des clients. Le coté client, doit être installé sur toutes les machines nécessitant d'accéder au serveur de BDD (un client peut éventuellement fonctionner sur le serveur lui-même) [5].

SQL Server :

SQL Server est un SGBDR (système de gestion de base de données relationnelle) entièrement intégré à Windows, ce qui autorise de nombreuses simplifications au niveau de l'administration, tout en offrant un maximum de possibilités. Il a une très grande capacité de gérer les données en conservant leur intégrité et leur cohérence. Il est chargé de stocker les données, vérifier les contraintes d'intégrité définies, garantir la cohérence des données stockées même en cas de panne et assurer les relations des données [6].

Firebird :

Firebird est système de base de données relationnel, comparable à des produits comme DB2 d'IBM, Oracle, SQL Server de Microsoft et le produit open source PostgreSQL. Le logiciel a deux principaux composants : le serveur de bases de données et l'interface applicative, communément appelée la « bibliothèque client ». La bibliothèque client est un composant nécessaire sur chaque station cliente dans le cadre d'un déploiement deux-tiers. Pour les déploiements multi-tiers, quand les utilisateurs accèdent aux bases de données à travers un middleware depuis un navigateur web ou autre "client léger", la bibliothèque cliente Firebird n'est pas déployée sur les stations des utilisateurs mais uniquement au sein du middleware [7].

III. Outils de manipulation de données

Le monde de l'informatique évolue très rapidement, alors que son but initial, était d'offrir des services satisfaisants, du point de vue vitesse d'exécution des tâches et obtention de statistiques plus précises dans le domaine des bases de données. Pour cela nous allons faire une étude détaillée de nombreux outils concernant :

III.1 Manipulation des données

La manipulation des données désigne toute chose concernant les données et leur traitement, navigation et même leur présentation par exemple l'exécution des requêtes, Auto-complétion, export/import etc.

Nous avons étudié certains outils qui sont comme suit :

1. MySQL Workbench**a. Définition :**

MySQL Workbench (Anciennement *MySQL administrator*) est un logiciel de gestion et d'administration de bases de données MySQL créé en 2004. Via une interface graphique intuitive, il permet, entre autres, de créer, modifier ou supprimer des tables, des comptes utilisateurs, et d'effectuer toutes les opérations inhérentes à la gestion d'une base de données. Pour ce faire, il doit être connecté à un serveur MySQL [8].

MySQL Workbench permet aux développeurs de créer un ou plusieurs modèles dans son interface et il offre différentes vues des objets qui sont en cours de conception (tables, vues, procédures enregistrées, etc.)

b. Fonctionnalités Côté Données:

MySQL Workbench visualise les données en deux manières (ResultGrid, Form Editor) qui contiennent de leur part des fonctionnalités comme effectuer une mise à jour sur une table, de lancer une recherche avancée (recherche par Id), de trier (selon un critère) et de filtrer toutes les données enregistrés dans une table, de copier et de coller les enregistrements sélectionnés. Concernant la création des requêtes il supporte l'Auto-Complétion et un détecteur d'erreur. Comme toutes les logiciels il fait l'exporter et l'import des données sous forme fichier CSV, SQL (seulement pour importer) ainsi il sauvegarde l'historique des requêtes SQL.

c. Avantages :

- MySQL Workbench est un logiciel rapide et robuste avec performances élevées en lecture ;
- Il fonctionne sur plusieurs plateformes : Windows, Linux, Mac OS X, et de nombreux types d'UNIX.
- Il a une bonne gestion de base de données, en plus il est multi-Database et multi-connexions.
- Il est un véritable outil ergonomique (Administration facile et rapide de MySQL).

d. Inconvénients :

- MySQL Workbench comporte moins de fonctionnalités par rapport à d'autres logiciels similaires ;
- Il se connecte à un seul SGBD MySQL.
- MySQL Workbench ne permet pas à l'utilisateur de choisir le séparateur d'importation et d'exportation des données CSV.

1. SQL Manager**a. Définition**

SQL Manager est un outil puissant d'administration et de développement de bases de données de serveur MySQL, PostgreSQL, SQLServer, Firebird, Oracle, Interbase, etc.

SQL Manager permet de Connecter à plusieurs hôtes/bases de données, de créer/modifier tous les objets de bases de données, de concevoir visuellement les bases de données, d'exécuter les scripts SQL, d'importer et d'exporter les données de bases de données, de gérer les utilisateurs et leurs privilèges et dispose d'autres services qui permettent de faciliter la gestion des bases de données [9].

b. Fonctionnalité Côté Données :

SQL Manager expose la visualisation des données en trois manières (GridView, FormView, PrintView). Nous remarquons que ces derniers contiennent plusieurs fonctionnalités par exemple : lancer une recherche avancée (recherche par Id), de grouper, de trier (selon un critère), de filtrer toutes les données stockées dans une table et de dupliquer des enregistrements sélectionnés. Et ce qui est logique de pouvoir exporter et importer les données sous plusieurs formats (CSV file, Texte file, PDF,...), en plus la possibilité d'exporter des données en tant que page PHP. Ensuite il donne la main d'exporter les données vers un script SQL en tant qu'instruction INSERT. En fin nous observons l'existence d'auto-complétion.

c. Avantages :

- SQL Manager Permet de se connecter à plusieurs serveurs MySQL, PostgreSQL, SQL Server, Firebird, Oracle, Interbase, etc. ;
- Il est Ergonomique et il contient plus de fonctionnalités par rapport aux autres logiciels similaires.

d. Inconvénients :

- Il existe une version pour chaque SGBD (SQL Manager for MySQL, for PostgreSQL, etc.). Par exemple, pour utiliser Oracle, il faut installer SQL Manager for Oracle, etc. ;

2. Heidi SQL :**a. Définition :**

Heidi SQL est un outil d'administration de base de données possédant un éditeur SQL. Son interface graphique permet de gérer des serveurs MySQL en manipulant les variables serveurs, les bases, table, vue, champ et données, import et export de données.

b. Fonctionnalités Côté Données :

Tout d'abord Heidi SQL accède à PostgreSQL, MySQL, SQL Server et de se connecter en ssh ou tcp/ip. Il duplique un enregistrement. Plus qu'il offre le service d'auto-complétion en gardant l'historique des requêtes. Il permet d'effectuer une recherche. Comme il fait l'import et l'export des données. Enfin il permet éditer et exécuter des scripts SQL.

c. Avantages :

- Coloration syntaxique des requêtes ;
- Instance multiple de connexions en l'ouvrant plusieurs fois.
- Enregistrement de profil de connexion.

d. Inconvénients :

- Il ne se connecte pas à tous les SGBD ;

3. IBExpert**a. Définition :**

IBExpert est une application cliente Interbase (créée par HK-Software) qui offre à l'utilisateur ou à l'administrateur un accès facile aux serveurs Interbase/Firebird ainsi qu'aux fichiers et objets Interbase/Firebird au travers d'une interface graphique. Existe en version complète (commercial) et personnelle (gratuit).

b. Fonctionnalité Coté Données :

En premier lieu il faut que nous validions tout type de modification sur les données soit par l'ajout, modification ou suppression, par conséquent il contient un détecteur d'erreur. En plus nous pouvons effectuer les opérations suivantes sur la colonne qu'on veut (sum, max, min, count et distinct etc.) via le composant show summary footer. Ensuite il met à jour une table et plus qu'il fait l'export et l'import des données sous forme fichier CSV.

c. Avantages :

- Il est ergonomique ;
- Nous pouvons annuler une transaction (rollback) après un commit par un Backup.
- Il peut se connecter à plusieurs BDD.

d. Inconvénients :

- Il ne se connecte pas à tous les SGBD seulement Firebird et Interbase ;
- Il n'est pas parmi les meilleurs pour éditer, importer/exporter.
- Il se bloque beaucoup.
- Puisque Firebird ne supporte pas les requêtes multi-bases. Donc impossible d'exécuter ce type de requêtes notamment avec IBExpert.
- Il n'est pas compatible avec Windows 10.

4. Toad for SQL Server**a. Définition :**

Toad for SQL Server est un client qui fonctionne sur les plateformes Windows. Il permet de gérer les dictionnaires de données, les tables, les index, etc. Il existe d'autres versions qui permettent d'accéder aux bases Oracle, PostgreSQL, MySQL, SQL Server, et DB2. En tant qu'administrateur de bases de données SQL, nous pouvons exécuter différentes fonctions en utilisant à la fois SQL Server Management Studio et Visual Studio. C'est pourquoi des fonctionnalités avancées d'administration et de réglage du code SQL permettant d'optimiser les performances pour rendre la tâche plus facile. Parmi ses fonctionnalités qu'il génère les requêtes. Plus il fait la comparaison et la synchronisation, aussi il sauvegarde l'historique des transactions. D'une part il améliore les performances des applications et l'automatisation des processus répétitifs. D'une autre part il fait une exécution groupée plus qu'il génère des documents et des rapports. Enfin il s'intéresse aussi de gérer la sécurité et d'importer une connexion a Toad for SQL Server [10].

b. Fonctionnalité côté données :

Tout d'abord il permet de se connecter à SQL Server. Ensuite il génère les requêtes en laissant les administrateurs de bases de données et les analystes exporter facilement des données vers leurs outils de prédilection. Plus l'accès rapidement aux données importantes (avec des fonctionnalités de rapports et de tableaux dynamiques croisés), ce qui permet de les analyser sur place et de les exporter vers une instance Excel en un seul

clic. Il fait la comparaison et la synchronisation en identifiant facilement les différences en comparant et en synchronisant les serveurs, les schémas et les données. Il fait la restauration des transactions dans le journal des transactions, sans devoir procéder à une restauration à partir d'une sauvegarde. Il améliore les performances des applications avec des fonctions automatisées d'optimisation et de réécriture des requêtes. Il est concerné aussi d'automatiser les processus répétitifs, notamment les comparaisons des schémas et des données. Il fait un rappel de scripts SQL déjà saisi en éliminant la nécessité de ressaisir les instructions SQL et T-SQL.

Assurez le suivi de toutes les instructions SQL exécutées au cours de votre session Toad.

Il donne aussi la chance d'exécuter en groupe (les scripts et des extraits de code sur plusieurs serveurs et instances). Il gère aussi la sécurité en créant, gérant et répliquant même des paramètres de sécurité pour tous les utilisateurs en élaborant et en exécutant des scripts de sécurité sur plusieurs serveurs. Il génère les rapports et les documentations en développant les rapports personnalisés à des fins d'administration et de développement, et exportant-les dans divers formats, notamment .XSL, .XML, .doc et .PDF. Il s'intéresse aussi à la recherche des objets par le biais de la localisation rapide du code spécifique pour évaluer au plus vite l'incidence d'un changement de nom ou de modifications du code en recherchant du texte dans des objets de base de données (par exemple dans les noms de colonne et le code SQL). Il gère en plus les projets car il enregistre et réutilise les puissantes fonctionnalités de gestion d'objets et de codage SQL et T-SQL. Il édite aussi T-SQL en réduisant le temps d'apprentissage du langage T-SQL, favorisant la cohérence et créant un code de haute qualité. Simplifiez le codage SQL avec des fonctionnalités avancées de finalisation de code, d'extraits de code et de rappel de scripts SQL. Il réduit le temps nécessaire à l'apprentissage des outils sur différentes plateformes, en profitant d'interfaces similaires tout en bénéficiant d'une expertise fonctionnelle approfondie adaptée aux nuances de chaque plateforme (Suite d'outils homogène). En plus il a comme fonctionnalité d'importer une connexion à Toad for SQL Server. Il a une gestion très riche des options [11].

c. Avantages :

- Administration centralisée ;
- Comparaison et synchronisation des objets et données SQL Server.
- Développements SQL et T-SQL simplifiés.

- Réglage automatisé des performances SQL.
- Prise en charge de SQL Server 2012 en plus des versions 2008, 2005 et 2003.
- Connexion à SQL Azure et utilisation des mêmes fonctionnalités Toad, comme pour une base de données SQL Server sur site.
- Nous pouvons changer la couleur de l'interface.

d. Inconvénients :

- Il n'est pas exécutable sur linux et Mac OS X ;

III.2 Data Profiling

Data Profiling est le processus qui consiste à récolter les données dans les différentes sources de données existantes (table, bases de données,...) et à collecter des statistiques et des informations sur ces données. Il est très proche de l'analyse des données.

Nous avons fait une étude de deux outils :

1. Talend Open Studio for Data Quality

a. Définition:

Talend Open Studio for Data Quality fournit des outils de développement et de gestion unifiés pour intégrer et traiter toutes les données dans un environnement graphique simple à utiliser. Dans le Studio, les utilisateurs peuvent accéder aux et examiner les données disponibles dans différentes sources de données et collecter des statistiques et des informations concernant ces données [12].

b. Fonctionnalités :

Il permet de sélectionner et d'afficher une table, d'exporter les données sous forme : fichier CSV, fichier Excel, fichier Html. Concernant le profilage des données il offre des graphes, des statistiques avec cinq types d'analyse :

- 1) l'**Analyse structurelle** : retourne une vue d'ensemble du contenu de la base de données. Elle calcule le nombre de tables et le nombre de lignes par table, pour chaque et/ou schéma, etc. Elle compte également le nombre d'index et de clés primaires, etc.
- 2) l'**Analyse inter-table** : Cette analyse explore différentes tables. L'analyse de redondance est un exemple dans lequel nous vérifions la relation entre deux tables, c.-à-d. cette analyse peut également être utilisée comme découverte de clé étrangère,

3) l'**Analyse de tables** : Retourne une vue d'ensemble du contenu de votre table. Elle calcule le nombre de lignes, le nombre de nuls etc. Elle compte également le nombre d'index et de clés primaires, etc.,

4) l'**Analyse de colonnes** : permet de profiler les données de colonnes,

5) l'**Analyse de corrélation** : est une analyse d'exploration. Elle permet d'explorer les relations entre les colonnes et peut aider à détecter des problèmes de qualité de donnée.

En plus il offre la personnalisation d'analyse, c.-à-d. il permet de sélectionner les champs de statistiques avant le profilage (Null count, unique count, blank count etc.), et de filtrer des données à profiler. Ensuite il permet d'analyser la qualité des données avec des graphiques, et d'afficher la liste des valeurs statistiques. Comme par exemple, pour les valeurs uniques, il affiche le nombre et pose la possibilité d'afficher des enregistrements uniques dans une grille.

c. Avantages :

- Permet de connecter à plusieurs sources de données (ex : MySQL, SQL Server, Oracle, PostgreSQL, SQLite,...) ;
- Permet d'établir plusieurs types d'analyse (structurelle, inter-tables, tables, colonne, corrélation).
- Plus de fonctionnalités par rapport les autres logiciels similaires.

d. Inconvénients :

- N'offre pas la possibilité de télécharger des rapports de Profilage de données ;
- Moins ergonomique par rapport à Toad Data Point.

2. Toad Data Point

a. Définition :

Toad Data Point est un outil de requête de base de données multiplateforme conçu pour quiconque a besoin d'accéder à des données provenant de sources multiples, de les intégrer rapidement et de les préparer, et de produire des ensembles de données et des rapports pour analyse. Il simplifie considérablement l'accès aux données, l'écriture de requêtes SQL et la préparation des données, à partir d'un seul outil.

b. Fonctionnalités :

Toad Data Point permet d'effectuer une mise à jour sur une table, il fait aussi le trie selon un critère et le filtrage de toutes les données enregistrées dans une table. Il permet la

création facile des requêtes, soit par la génération d'un simple script SQL (select, insert...) ou bien grâce à une interface graphique (seulement pour select). L'utilisateur peut effectuer un export de données sous forme : fichier CSV, fichier Excel, fichier Html.

En arrivant au dernier point c'est le profilage des données (des graphes, des statistique) suivis par la génération de leur rapport.

c. Avantages :

- Toad Data Point Permet de connecter à plus de 50 sources de données (relationnelles, non relationnelles, locales ou hébergées dans le Cloud) à partir d'un seul outil ;
- Il est ergonomique et rapide.
- Il facilite le profilage et le nettoyage des données.

d. Inconvénients :

- Toad Data Point ne donne pas la main à l'utilisateur pour choisir l'emplacement d'exporter les données d'une table ;
- Toad Data Point permet de générer des rapports de profilage des données seulement pour les tables.
- Les rapports générés ne sont pas dynamiques (navigable).

III.3 Génération de données

C'est la création d'un ensemble de données de test nécessaires selon le besoin de l'utilisateur. Qui consiste à remplir les tables dans une base de données, puis la spécification des détails concernant la génération des données pour des tables et des colonnes spécifiques c.-à-d. pour chaque colonne, il faut avoir un générateur de données qui la produit d'un type particulier.

La génération de données inclut plusieurs générateurs de données intégrés pour la génération de différents genres de données. Par exemple, le générateur des valeurs aléatoires simple ou avec masque, séquentiels ou valeurs Constante, valeurs Nul, et via dictionnaires, etc.

Un générateur de données de test est couramment utilisé pour les bases de données de test et logiciel de gestion de base de données (SGBD).

Lorsque nous exécutons des tests unitaires de base de données, nous pouvons spécifier un autre plan de génération de données pour chaque projet de test. Par conséquent, nous pouvons initialiser la base de données à un état différent pour chaque groupe de tests.

Nous avons fait l'étude de :

1. DTM Data Generator :

a. Définition :

DTM Data Generator est un produit logiciel qui produit des lignes de données et des objets de schéma à des fins de test : la population de la base de données d'essai, l'analyse de la performance, l'épreuve d'assurance qualité ou l'exécution des tests de chargement. Le générateur a été conçu pour fournir aux développeurs et aux ingénieurs d'assurance qualité des tableaux de tests réalistes et de haute qualité. Il crée automatiquement des valeurs de données et des objets de schéma optionnels (tables, vues, procédures, déclencheurs, etc. [13].

b. Fonctionnalités:

Il permet de se connecter à SQL Server, MySQL, Interbase/Firebird, Oracle, IBM DB2, PostgreSQL et de tester si la connexion est valide. Plus qu'il offre la gestion d'une table (Ajouter, modifier et supprimer). Il permet aussi d'exécuter des requêtes SQL et de faire un import/export.

L'Assistant Règle de ce dernier permet aux utilisateurs de créer un projet de génération de données via quelques clics. L'analyseur du schéma intelligent rend les données réalistes sans modifications de projet supplémentaires et la fonctionnalité « données par exemple » rend les données plus réalistes sans efforts supplémentaires.

Ensuite le module de vérification du schéma analyse la structure de la base de données cible avant chaque exécution pour empêcher le remplissage des tables modifiées. La fenêtre d'aperçu montre des données d'échantillon à générer en un clic et permet de gagner du temps dans le temps de conception des règles. Aussi Le générateur de données de test peut non seulement préparer des données de test : des tables, des vues et des procédures peuvent également être générées en vrac.

Le rapport de projet (échantillon) permet d'examiner les règles de génération de données. Le rapport de population de base de données permet aux utilisateurs d'analyser les résultats d'exécution (échantillon). BLOB loader offre un transfert massif de données binaires vers la base de données.

Le SDK de génération de données est un moyen simple caractérisé par l'ajout des fonctionnalités qui font la création des données de test à votre propre application ou site Web.

c. Avantages:

- Permet de gérer les règles de génération ;
- Il est très riche.

d. Inconvénients :

- Difficile à utiliser ;
- Il ne se connecte pas à quelque SGBD.

2. Generate Data :

a. Définition :

Un générateur de données de test est un outil de logiciel spécialisé qui génère des données fausses ou simulées pour une utilisation dans les applications logicielles de test car ils nécessitent généralement de grandes quantités de données afin de faire des tests.

Les données générées peuvent être aléatoire ou spécifiquement choisies pour créer un résultat souhaité [14].

b. Fonctionnalités :

Il permet de faire tout type d'export (CSV [Windows,MAC,Unix] EXCEL HTML ...). Nous pouvons appliquer ajout ou suppression d'un plugin, en réinitialisant les plugins. Cela mettra à jour la base de données et nous assurons l'accès aux plugins dont nous avons besoin. Il offre le partage des données avec d'autres personnes. Comme il sert à ouvrir les données générées dans une fenêtre/onglet ou dans la page elle-même ou bien les télécharger dans un fichier normal où .ZIP. En plus il est caractérisé par sa connexion à Oracle, PostgreSQL, SQL Server, MySQL. Ensuite il génère des données avec insert, drop, create, update. Il contrôle le nombre des lignes. Il assure la protection des noms des tables et des champs et le sauvegarde de l'historique. Il y a l'ajout d'une colonne d'auto-incrémentation. Il offre deux formats de la structure des données complexe/simple. Il fait un export sous plusieurs langages de programmation (perl, Ruby, JavaScript, PHP). Plus la génération des données de test aléatoire.

c. Avantages :

- Facile à utiliser ;
- L'existence d'auto-incrémentation.
- Utilise des dictionnaires.

d. Inconvénients :

- N'importe pas le schéma à partir d'une BDD, il faut spécifier les champs de la table sur l'application ;
- Génération limitée de données.

3. Mockaroo**a. Définition :**

Mockaroo est un générateur de données de test en ligne gratuit qui permet de créer facilement des données réalistes pour tester une application, ou pour remplir des données de test dans une base de données. C'est un très bon produit et est parfait s'il y a un besoin pour générer des données d'échantillons grands, réalistes rapidement. C'est un site en ligne [15].

b. Fonctionnalités :

Il est caractérisé par une interface graphique simple et ergonomique. Il génère 1000 lignes. Il expose aussi tous types d'export : CSV, JSON, SQL, et Excel et sous la machine Windows/Linux. L'importation des données csv/Excel/SQL. En plus il clone le schéma.

c. Avantages :

- Facile à utiliser ;
- Importer un schéma.

d. Inconvénients :

- Il ne se connecte pas à une BDD pour l'import du schéma et la génération des données ;
- Il génère que 1000 données dans la version gratuite et au-delà il faut passer à la version payante.
- Il ne se connecte à aucun SGBD.

4. Spawner**a. Définition :**

Spawner est un programme qui peut s'exécuter sur l'hôte serveur et écouter les requêtes. Cette application contient un module appelé géniteur qui fonctionne comme suit :

- Il accède au fichier de configuration des métadonnées pour obtenir des informations sur la connexion au Metadata Server.

- Il se connecte au serveur de métadonnées pour les informations de configuration.

b. Fonctionnalités :

Ce logiciel offre une génération illimitée de données en proposant plusieurs types concernant les champs par exemple : Integer, Varchar, FullName, FirstName etc., en fait les enregistre sous fichier SQL, en faisant des inserts, insert ignore ou replace dans la table que nous voulons peupler. Il donne la main d'exécuter les fichiers SQL télécharger dans n'importe quelle SGBD. Il génère aussi des tests data dans une BDD spécifiée ou en utilisant un INSERT STATEMENT, ou bien en créant un fichier Text/XML. Ensuite il permet de dupliquer les tables, désactiver/activer une table et exécuter des requêtes SQL via SQL Console en gardant l'historique.

c. Avantages :

- Il est ergonomique ;
- Nombre de données à générer non limité.
- Il peut être installé sur plusieurs machines : (z/OS, OpenVMSAlpha, UNIX (AIX 64, HP-UX, IPF, HP 64, Tru64 UNIX, Solaris 64, RedHat Linux on Intel), Windows NT/XP/2000).
- Il offre l'utilisation des mécanismes de sécurité de serveur normal pour protéger les données sensibles.

d. Inconvénients :

- Il ne se connecte pas aux SGBD.

5. Yan Data Ellan**a. Définition :**

Yan Data Ellan est un site gratuit qui peut générer des données de test dans plusieurs domaines dans plusieurs formats.

b. Fonctionnalités :

Il offre le choix d'un SGBD (Oracle, SQLite, MySQL, PostgreSQL, SQL Server). Après la connexion nous aurons la chance de générer jusqu'à 10 000 lignes de données et les exporter sous plusieurs formats de fichiers (CSV, Excel, SQL, JSON, HTML et XML). aussi nous pouvons faire des insert/update. Nous pouvons de plus affecter un nom à la table/fichier qui contient les données et d'établir un auto-incrément à un champ. Il est riche car il génère des données réelles, valides qui peuvent être

spécifiées dans différents pays : Australie, Canada, France, Allemagne, Italie, Espagne etc. Il contient des data types variés (personne données, données de localisation, données de l'entreprise, données de banque, données mathématiques, données numériques, données –IT....).

c. Avantages :

- Facile à utiliser ;

d. Inconvénients :

- Génération limitée de données (que 10000 lignes) ;
- Il ne se connecte pas aux SGBD.

III.4 Etude Comparative

Table 1. Tableau comparatif de manipulation des données :

	Payant	Open source	Ergonomie	Multi-connexion	Auto-complétion	Créateur	SGBD supportés	SE	Licence
SQL Manager	Oui (depuis €95.00)	Non	A++	Non	Oui	EMS Database Management Solutions	MySQL, PostgreSQL, SQL Server, Firebird, Oracle, Interbase, ...	Windows XP/Vista/7/8/10	Propriétaire
MySQL Workbench	Non	Oui	A	Oui	Oui	Oracle Corporation	MySQL	Windows Linux et Mac OS	Community Ed: GPL Standard Ed: Commercial Proprietary
IBExpert	Oui (3500\$)	Non	A+	Non	Oui	?	Interbase Firebird	windows 7, linux	?
Heidi SQL	Non	Oui	A+	Oui	Oui	?	MySQL, PostgreSQL, SQL Server	Windows (7/10) ubuntu	GPL
Toad for SQL Server	Oui (gratuit 60 jours) (701.00 \$)	Non	B	Oui	Oui	Quest Software	SQL server	Windows, Mac, Unix	propriétaire

Les outils de manipulation de données fournissent une interface de connexion à plusieurs SGBD, chaque outil supporte un nombre spécifié de SGBD. Nous avons étudié quelque outil selon certaines fonctionnalités et critères. Le tableau ci-dessus dresse le résultat de notre comparaison.

La plupart des outils qui sont performants et qui présentent un grand nombre de fonctionnalités nécessaires sont payant tel que « SQL manager » et « Toad for SQL Server ». Mais malgré leur puissance en termes de fonctionnalités, ils possèdent quelque inconvénient. Nous remarquons que la majorité des outils proposent une interface graphique ergonomique à part « Toad for SQL Server ». Tous les outils offrent le service d'auto-complétion. A la fin, Concernant la compatibilité entre environnements nous constatons que « MySQLWorkbench » et « Toad for SQL Server » sont plus compatible avec d'autre environnement par rapport à « IBExpert » et « SQLManager ».

Table 2. Tableau comparatif de Data Profiling :

	Payant	Open Source	Ergo-Nomie	Multi-DataBase	Auto-Complétion	Créateur	SGBD supportés	Système D'exploitation	Licence	Rapport navigable	Rapport simple
Talend Open Studio for data quality	Gratuit	Oui	A+	Oui	Non	Talend	MySQL, PostgreSQL, SQL Server, Oracle,...	Window Linux/ Solaris/ Mac Os	Apache	Non	oui
Toad Data Point	Oui (€135.94)	Non	A+	Oui	Oui	Quest Software	50 sources de données (MySQL, PostgreSQL, Oracle,...)	Windows 7/8/8.1/10	Propriétaire	Non	oui

Les outils de Data Profiling fournissent une interface de connexion à plusieurs SGBD, comme nous observons dans le tableau ci-dessus les deux offrent la possibilité de se connecter à MySQL, PostgreSQL, Oracle etc. Puisque Toad Data Point est plus riche en termes de fonctionnalités, il est payant. Par contre « Talend Open Studio » est gratuit. Nous remarquons que les deux proposent une interface ergonomique. Nous notons que « Toad Data Point » est plus compatible avec d'autre environnement par rapport à « Talend Open Studio » qui fonctionne qu'avec les versions de Windows. Nous observons aussi qu'ils portent des rapports simple, par contre ils n'y a pas des rapports navigable.

Table 3. Tableau comparatif de Génération de données:

	Payant	Open Source	Ergono- Mique	SGBD support	Génération de clés étrangères	Système D'exploitation	Licence
DTM Data Generator	Oui US \$1149	Non	A+	Oracle ,MySQL , Firebird, IBM DB2, SQL Server, PostgreSQL	Oui	Window xp,10,8,	Demo
Generate Data	Oui \$ 799	Non	A+	Oracle, PostgreSQL , SQL Server, MySQL	Non	Windows ,Unix	GNU GPL
Mockaroo	Non	Oui	A+	Aucun	Non	Window/Linux	?
Spawner	Oui \$100	Non	A	Aucun	Non	Window 7	GNU General Public
Yan Data Ellan	Non	Oui	A+	Aucun	Non	Window 7	?

D'après le tableau ci-dessus nous observons que tous les outils de génération offrent une interface graphique ergonomique. Notons que « DTM Data Generator » et « Generate Data » supportent la connexion à plusieurs SGBD par rapport aux autres qui n'offrent aucun accès. Nous remarquons aussi que « DTM Data Generator » est cher par rapport à « Generate Data » et « Spawner » parce qu'il offre plus de fonctionnalités par rapport à l'autre, par contre le reste des outils proposent des versions gratuites « open source ».

Comme ils sont compatibles avec plusieurs environnements. Concernant la génération des clés étrangères, il y a un seul logiciel qui traite ce cas c'est « DTM Data Generate»

III.5 Conclusion

En constat, d'après l'étude présentée dans le chapitre 1, un besoin d'une solution de réaliser un module de manipulation de données plus efficace, rapide et fiable qui englobe tous les fonctionnalités existant ou inexistant dans les autres logiciel concurrents en traitant tous les besoins et les appliquer dans notre application. Par conséquent, il sera très intéressant, car c'est une chance et un honneur à nous de créer un module déjà plus efficace et fiable surtout si nous prendrons en considération les logiciels concurrents.

CHAPITRE II

ANALYSE ET CONCEPTION

I. Introduction

Après avoir définie la problématique et les besoins dans l'introduction générale et après avoir effectué une étude détaillée sur les solutions d'administration et de manipulation de données existant sur le marché, nous allons entamer la phase d'analyse et conception de notre application. Ce chapitre va définir le processus de développement que nous avons suivi au long de notre projet et la conception et la modélisation de notre système. Nous allons décrire la façon dont le système va fonctionner en lui donnant une forme et une architecture.

Pour la conception de notre projet, nous avons choisi d'utiliser le langage UML pour détailler les aspects fonctionnels, dynamiques et structurels.

Nous avons essayé de résoudre les problématiques déjà citées et proposer une application plus efficace et puissante par rapport aux autres applications déjà présentées.

II. Processus de développement

La méthode Scrum se range sous la bannière d'un mouvement, l'agilité. Elle possède des valeurs et des principes et se met en œuvre avec des pratiques. De ce mouvement novateur émergent les méthodes agiles dont Scrum est actuellement la plus populaire [16].

Scrum est une méthode pour gérer les projets de façon agile. Nous l'avons choisie comme méthode de gestion de notre projet pour développer la nouvelle version « Module de manipulation de données », parce qu'elle permet d'offrir une meilleure visibilité, une forte inspection et une meilleure adaptation par rapport à ce que d'autres méthodes puissent offrir, à l'exemple du modèle en spirale qui a été utilisé pour le développement de la version précédente [17].

Scrum utilise un principe de développement itératif qui consiste à découper le projet en plusieurs étapes que nous appelons « itérations » ou « sprints ». Ces itérations sont constituées d'un ensemble de sous-besoins appelé « User story », en détaillant les différentes fonctionnalités qui seront développées. Un planning correspondant aux tâches nécessaires pour le développement de ces fonctionnalités sera établi.

Scrum se base sur une équipe avec différents rôles : « Product Owner », « Scrum Master » et « Team ».

- Product Owner : il s'agit d'avoir une vision de ce que nous voulons construire et de transmettre cette vision à l'équipe de développement. Le Product Owner définit le produit et priorise les fonctionnalités voulues.
- Scrum Master : il est chargé d'assister et guider l'équipe de développement pendant le sprint (itération) et facilite la résolution des problèmes.
- Team (équipe de développement) : elle a l'autonomie pour déterminer quand et comment achever ses travaux. Les membres travaillent dans une seule et même pièce. L'équipe livre un produit utilisable à la fin de chaque sprint.

Pour permettre de suivre les tâches à réaliser durant un sprint, nous utilisons un graphique « Burn Down Chart ». Il possède en abscisse le temps et en ordonnée le nombre de Taches restantes. La figure II.1 montre le Burn Down du premier sprint. Ce sprint a commencé le 30-01-2017 jusqu' à 25-02-2017 avec 20 tâches. La ligne grise montre un déroulement parfait du projet. La ligne bleue montre les dates butoirs définis au début du sprint. La ligne orange montre les dates de fin réel de chaque tâche.

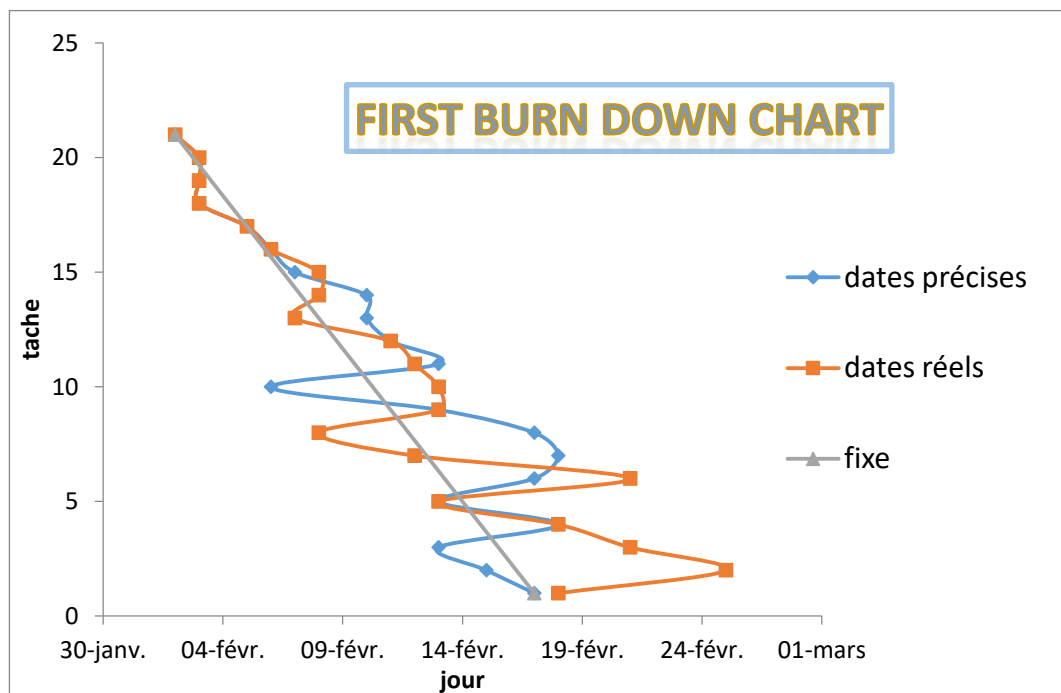


Figure II.1: Burn Down Chart: Iteration 1.

III. Sprint

Un sprint est le terme utilisé dans SCRUM pour itération. Un sprint est une itération qui génère un nouvel incrément (incrémental) et peut aussi enrichir un incrément obtenu dans un sprint précédent (itératif) [18]. Nous avons réalisé quatre itérations, nous avons donné un nom pour chacune. Durant chaque sprint, nous avons tombé les bugs que nous avons traités. A la fin de chaque sprint, après la réalisation de toutes ses fonctionnalités, nous passons à la phase test.

1. Itération 1 :

« بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ » est la première itération de notre application. Elle a duré presque trois semaines « 30-01-2017 au 25-02-2017 »

Nous citons ci-dessous quelques fonctionnalités de cette itération :

- La gestion des connexions dans SGBD;
- La mise à jour d'une table.
- L'exécution d'une requête select.
- L'export en CSV d'une table.

2. Itération 2 :

« Naissance » est la deuxième itération de notre application. Elle a duré deux semaines du « 26-02-2017 au 12-03-2017 ».

Nous citons ci-dessous quelques fonctionnalités de cette itération :

- L'implémentation du FormView (vue en formulaire) ;
- La réalisation de la duplication.
- L'implémentation de la recherche multicritères.
- L'export des données vers un script SQL en tant qu'instruction INSERT.
- La gestion des multithreads des requêtes.
- L'implémentation du PrintView (rapport PDF imprimable).

3. Itération 3 :

« Travail dur » est la troisième itération de notre application. Elle a duré 2 mois du « 19-03-2017 au 10-05-2017 ».

Nous citons ci-dessous quelques fonctionnalités de cette itération:

- La génération d'un script d'une table en tant que select, delete, insert;
- L'utilisation du XML plus export en XML.
- La génération des données.

- L'implémentation de l'auto-complétion avec un automate d'état finis.

4. Itération 4 :

« La fin » est la quatrième itération de notre application. Elle a duré 25 jours du « 12-05-2017 au 5-06-2017 ».

Nous citons ci-dessous quelques fonctionnalités de cette itération :

- La réalisation de Data Profiling;
- La refonte de l'IHM de toute l'application.
- Génération du data profiling d'une base de données et d'une table.
- La présentation et l'explication du code source de notre application aux développeurs de l'entreprise Soft Builder.

IV. Architecture

Après l'étude qu'on a faite pour réaliser ce module et dans le but d'avoir le meilleur code source et un développement plus efficace et rapide, nous avons suivi dans notre travail cette architecture montrée dans la figure ci-dessous. Tout d'abord nous avons commencé par le point commun et le plus important qui est la gestion des connexions aux cinq SGBD gérés par notre application (MySQL, PostgreSQL, Oracle, SQL Server et Firebird). Notre système se compose en trois parties. Dans la première nous avons le Data Profiling qui contient les fonctionnalités de calcul et génération des rapports sur les statistiques lié aux données). La deuxième partie concerne la génération des données qui se divise elle-même en trois parties : définition de la méthode de génération, la génération de données et l'export de ces données dans un fichier ou leurs introductions dans une BDD. Après nous avons la manipulation des données qui est scindé en cinq fonctionnalités riche : export sous plusieurs formats (CSV, SQL, XML), nous avons aussi l'import sous formats CSV et SQL, traitement des données, auto complétion et impression sous format PDF.

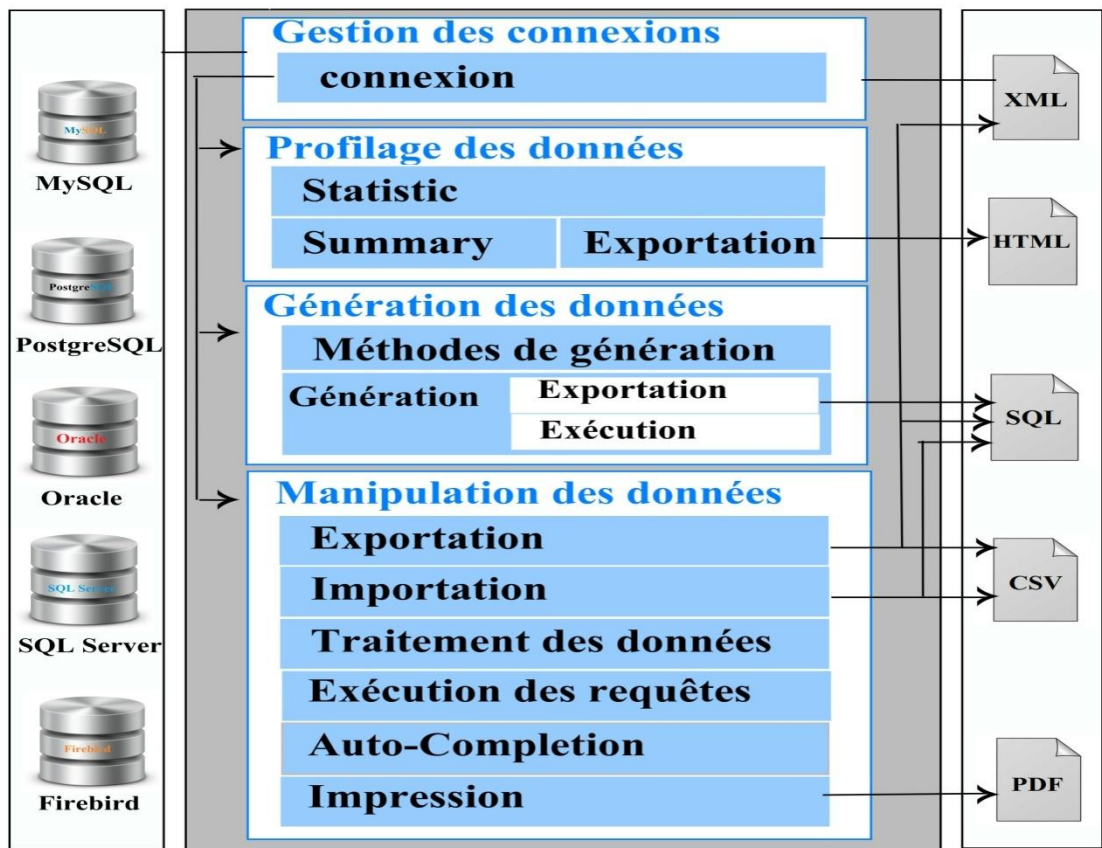


Figure II.2: Architecture.

V. Conception avec UML

UML, c'est l'acronyme anglais pour « Unified Modeling Language ». Qui est traduit par « Langage de modélisation unifié ». La notation UML est un langage visuel constitué d'un ensemble de schémas, appelés des 'diagrammes', qui donnent chacun une vision différente du projet à traiter [19].

Dans ce qui suit, nous allons faire la modélisation de l'application. Pour cela, nous avons utilisé les trois diagrammes d'UML suivants :

- Diagramme de cas d'utilisation.
- Diagramme de séquence.
- Diagramme de classe.

V.1 Diagramme de cas d'utilisation

Un cas d'utilisation (Use Cases) est un diagramme qui modélise une interaction entre le système informatique à développer et des acteurs interagissant avec le système.

Il permet aussi de définir les besoins des utilisateurs et les fonctionnalités du système,

Les éléments de base de cas d'utilisation sont :

Acteur :

Entité externe (humain, matériel,...) qui interagit avec le système étudié en échangeant de l'information. Notre système comporte un seul acteur principal qui est l'utilisateur du système et cinq acteurs secondaires qui sont : BDD, Manipulateur XML, Manipulateur CSV, Manipulateur SQL, Manipulateur PDF, Manipulateur HTML.

Cas d'Utilisation :

Un cas d'utilisation représente une fonctionnalité fournie par le système, en répondant à une action d'un acteur.

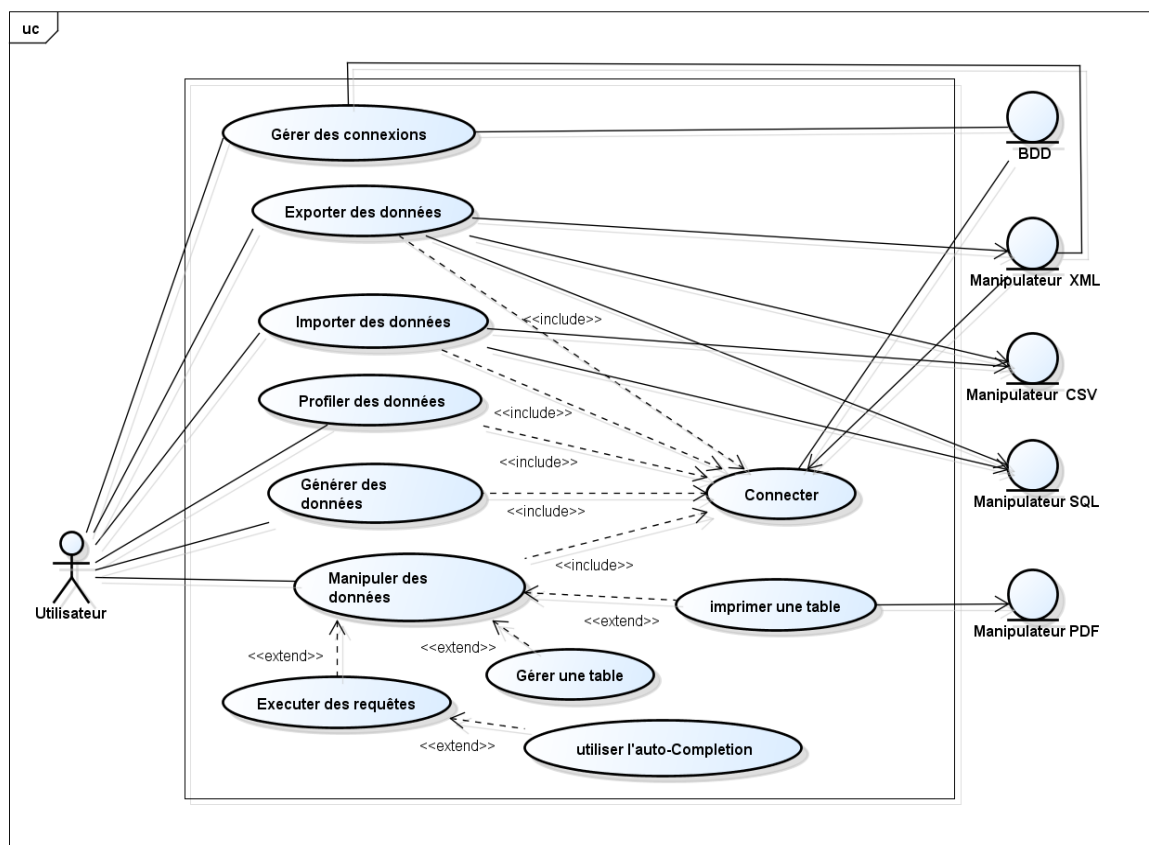


Figure II.3: Diagramme de cas d'utilisation.

Chaque usage que l'acteur fait du système est représenté par un cas d'utilisation. Ce dernier décrit le comportement du système du point de vue de son utilisateur.

Le diagramme de cas d'utilisation ci-dessus définit l'acteur principal qui interagit avec le système, en déterminant les besoins de l'utilisateur et tout ce que doit faire le système pour son acteur. L'utilisateur a la possibilité de gérer la connexion, faire export/import des données sous format SQL, XML, CSV, ainsi peut faire un data profiling, générer les données et manipuler les données en exécutant des requêtes, gérant une table et imprimer une version sous format PDF et tout ça se fait par l'obligation d'être connecté à n'importe quel SGBD qu'on veut utiliser.

V.2 Diagramme de séquence

Le diagramme de séquence est une représentation graphique des interactions entre des interactions entre les acteurs et le système.

- **Diagramme de séquence « connexion » :**

L'utilisateur doit lancer l'application. Puis il doit cliquer sur le bouton connecter, la connexion sera vérifiée par le fichier XML de gestion de connexion en vérifiant le nom et les autres informations et enfin soit sa connexion est bien réussite et créé soit le système affichera un message d'erreur.

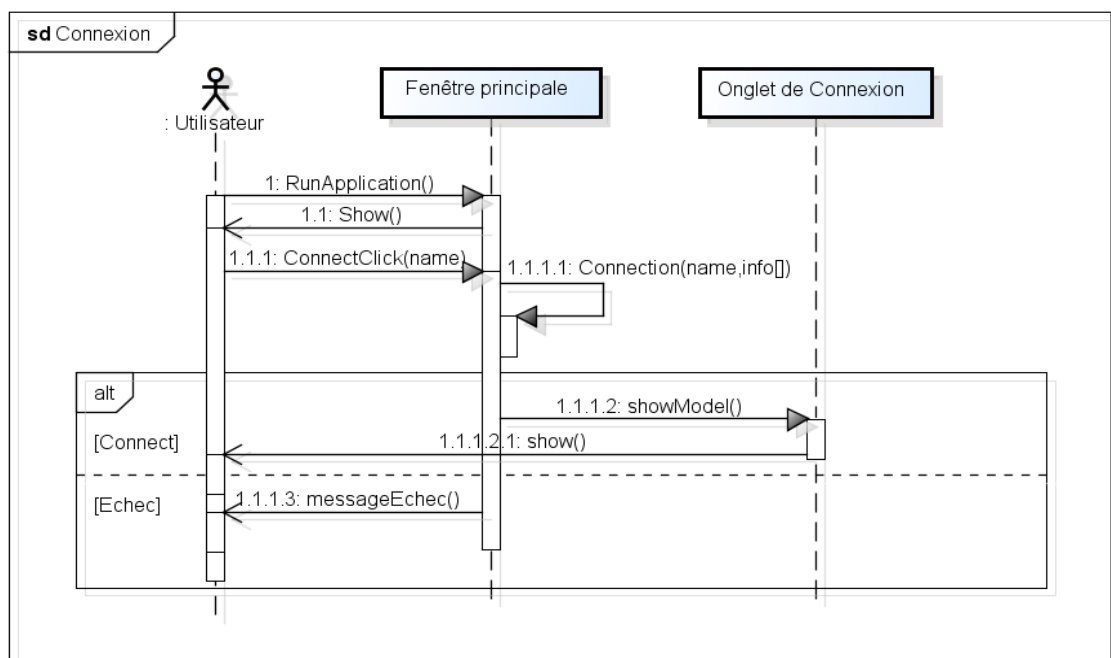


Figure II.4: Diagramme de séquence « Connexion ».

- **Diagramme de séquence « gestion des connexions » :**

L'utilisateur a le droit de gérer ses connexions. Il peut effectuer un ajout donc il doit remplir un formulaire pour ajouter une nouvelle connexion, et il doit cliquer sur ajouter et à partir de cela le programme doit vérifier automatiquement si les informations de connexion (User, password, port,...) remplis correspondent aux informations lié à la base de donnée, si oui la connexion sera ajouter sinon il affichera un message d'erreur. Il peut aussi supprimer une connexion en cliquant sur le bouton supprimer. Il peut regrouper les connexions par : nom d'utilisateur, SGBD, nom de base de données, défaut.

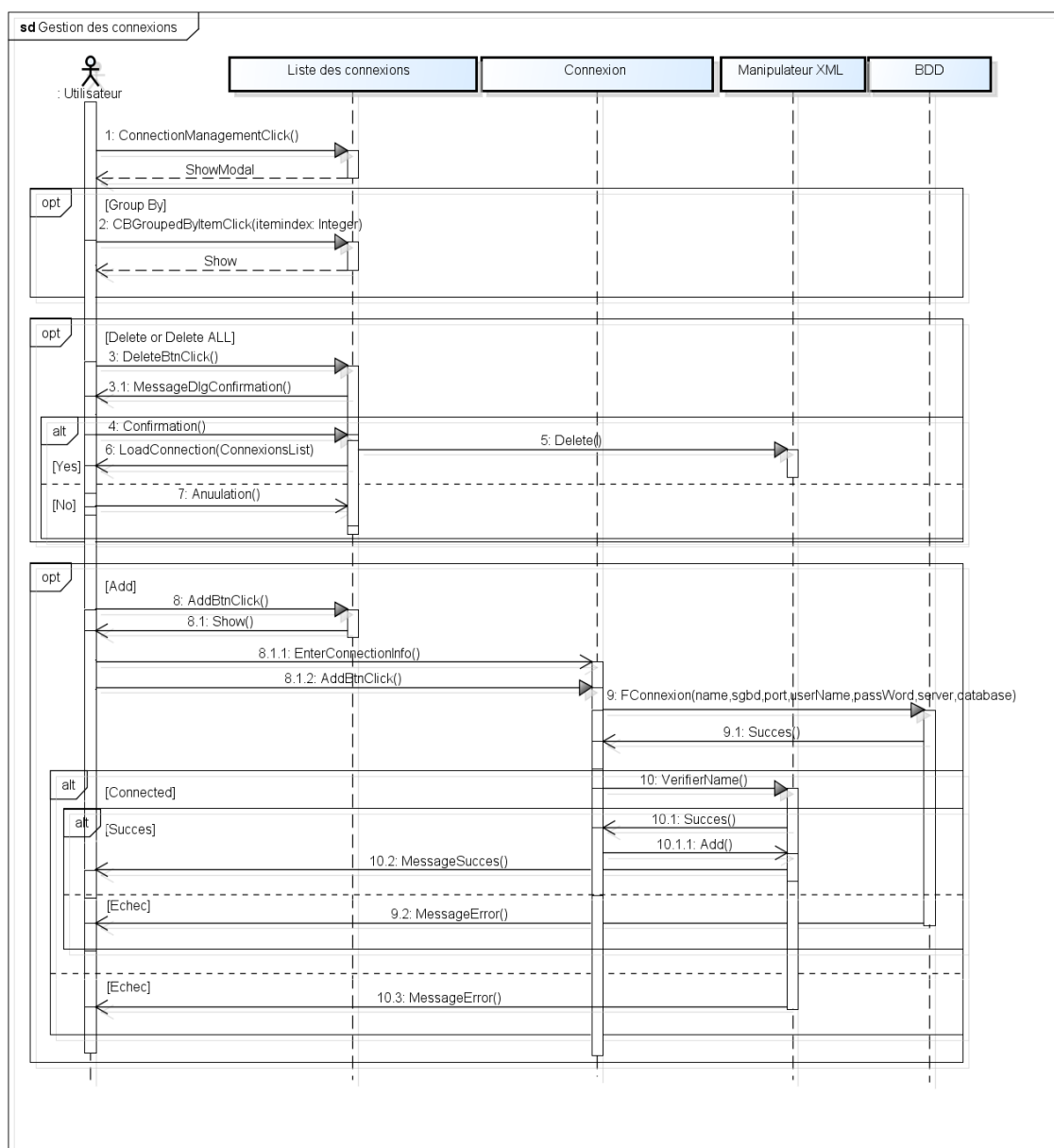


Figure II.5: Diagramme de séquence «gestion des Connexions».

- **Diagramme de séquence « manipuler les données » :**

L'utilisateur doit être connecté pour qu'il puisse gérer sa table. Tous d'abord il identifie sa table. Il clique sur GridView et de cela il peut faire un tri ascendant ou descendant sur n'importe quel champ de la table.

Il peut aussi exécuter des requêtes SQL, quand cette dernière est exécutée elle sera sauvegardée dans un historique, puis le système affichera le résultat de la requête si la requête est bien écrite sinon il affichera un message d'erreur. Comme il peut faire un export en CSV en cliquant sur Export CSV, ensuite il doit sélectionner les champs à exporter, en plus l'utilisateur doit choisir un séparateur de données, puis il doit cliquer sur export en choisissant aussi le chemin.

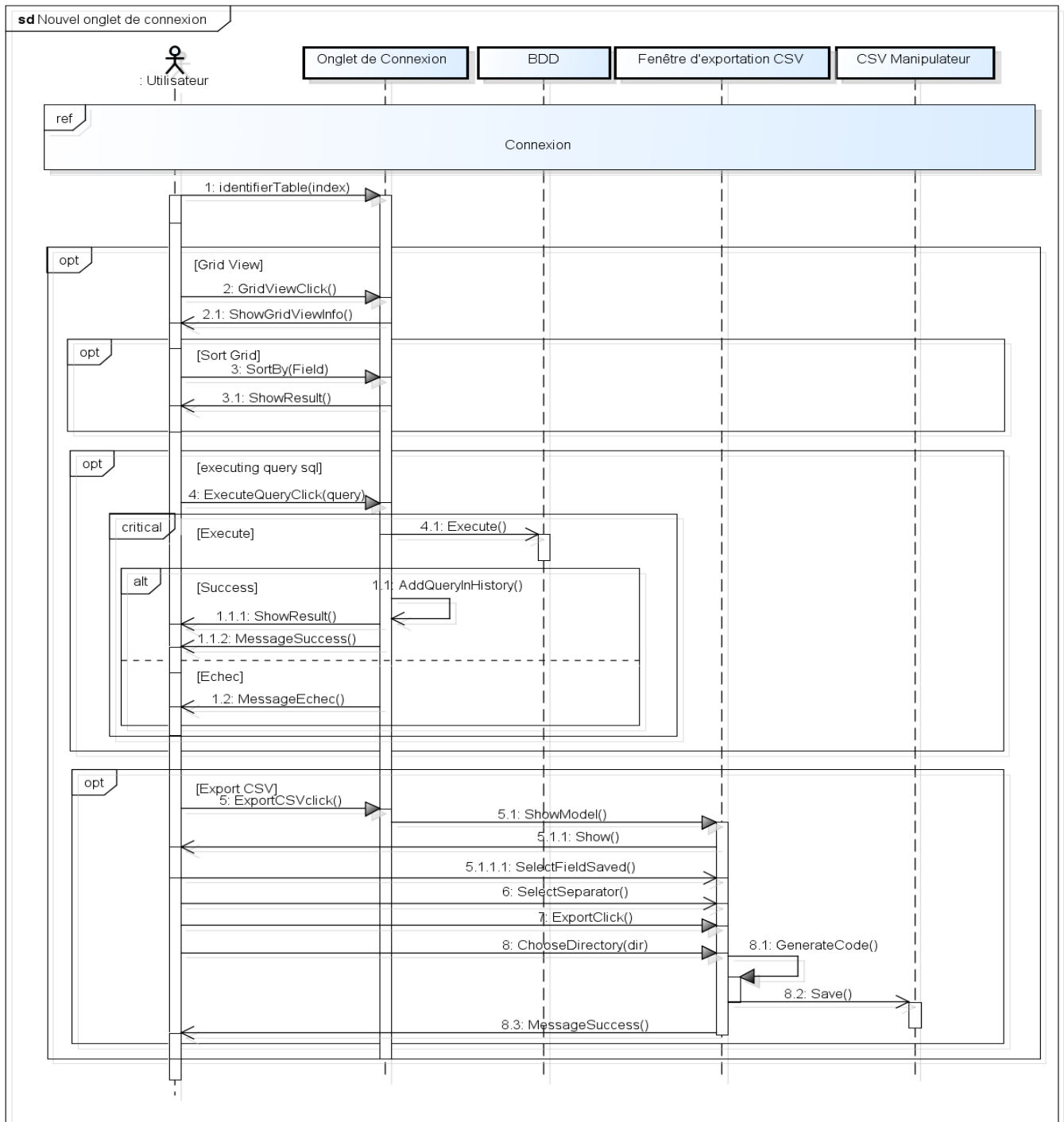


Figure II.6: Diagramme de séquence «Manipulation des données».

- Diagramme de séquence « Générer les données » :

L'utilisateur peut générer des données. Il doit d'abord identifier sa table, ensuite il doit entrer le nombre de lignes à générer selon son besoin, après il doit choisir les types pour chaque champ qui fait partie de la table identifiée, puis il doit cliquer sur le bouton générer et là il peut choisir de le sauvegarder sous format fichier en faisant appel à SQL Manipulateur où il exécute directement sa génération de données. Ces deux alternatives se terminent soit par succès soit par un échec.

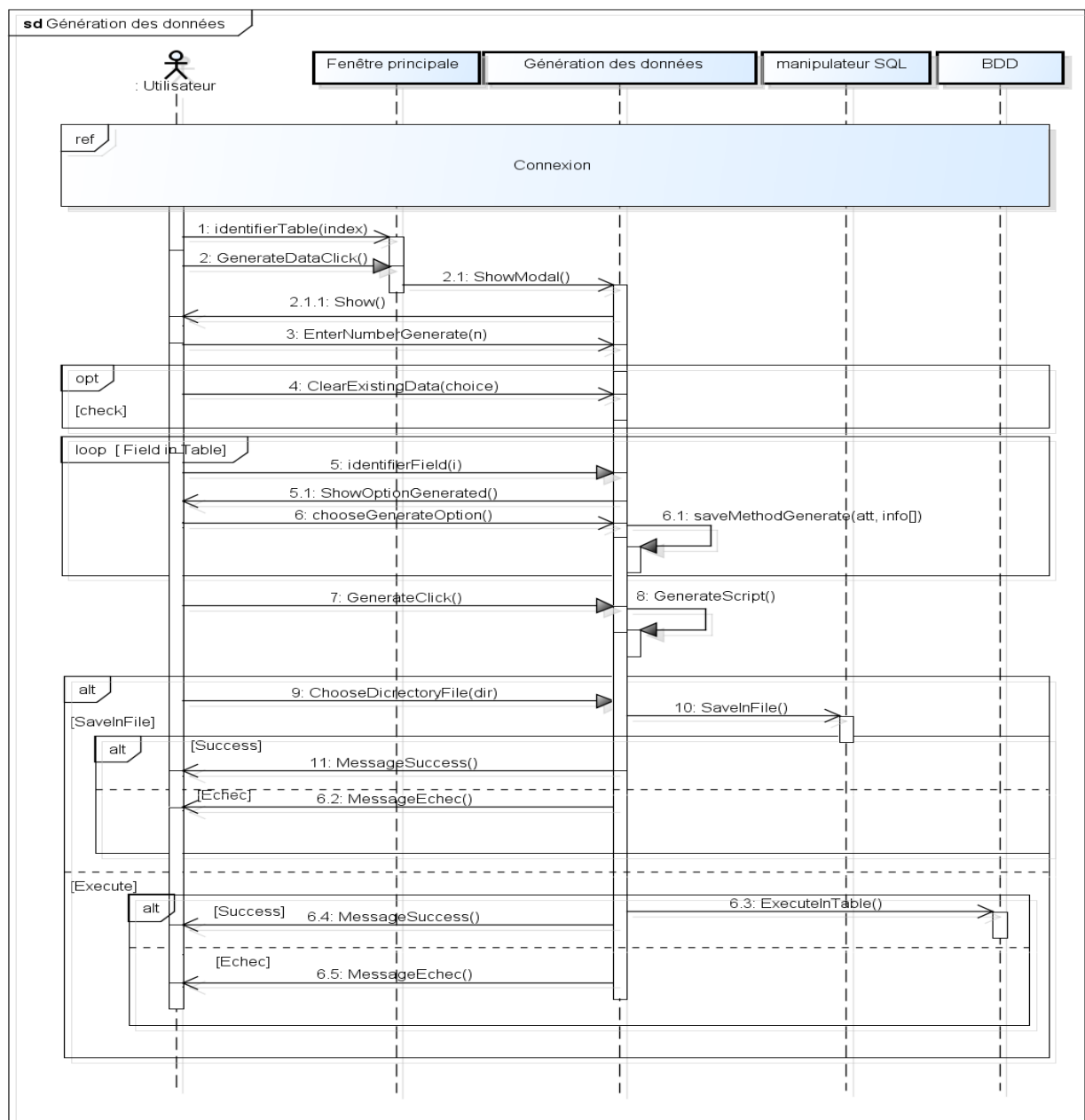


Figure II.7: Diagramme de séquence «Génération des données».

- **Diagramme de séquence « Data Profiling » :**

L'utilisateur peut générer des rapports sur les données de sa BDD. Tout d'abord il doit identifier sa table et il doit cliquer sur data profiling. Il peut choisir summary (résumé des statistiques), puis il génère un rapport web et à partir de cela le code HTML, CSS, JavaScript et BootStrap va être généré et enregistré et enfin notre système affichera un message d'erreur en cas d'échec sinon l'opération se terminera avec succès. En plus de la vue summary l'utilisateur peut avoir des statistiques détaillées.

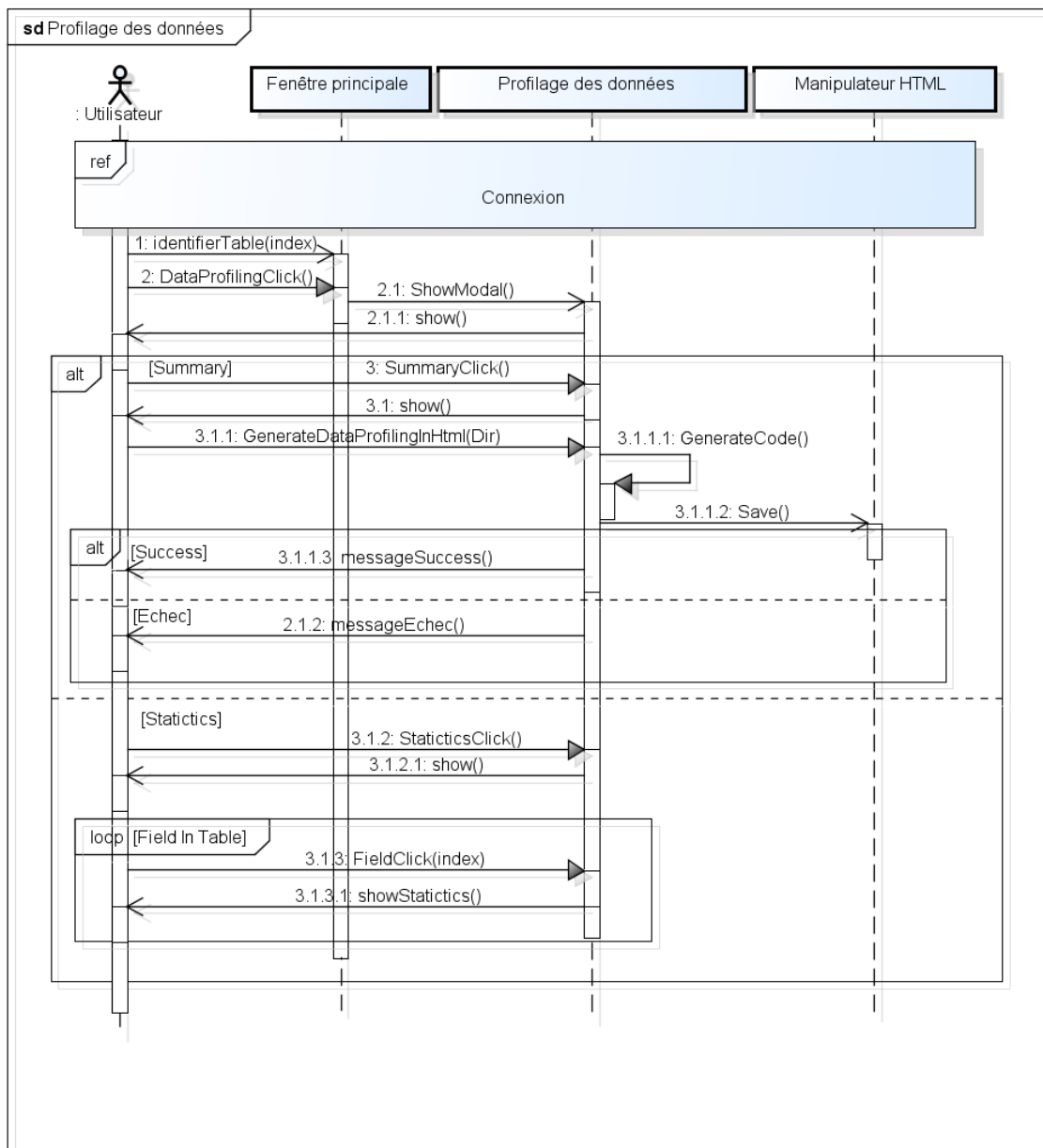


Figure II.8: Diagramme de séquence «Data Profiling».

V.3 Diagramme de classe

Un diagramme de classe permet de représenter les classes d'un système et les différentes relations entre celles-ci, c.-à-d. une représentation statique.

La figure suivante représente le diagramme de classe de notre application :

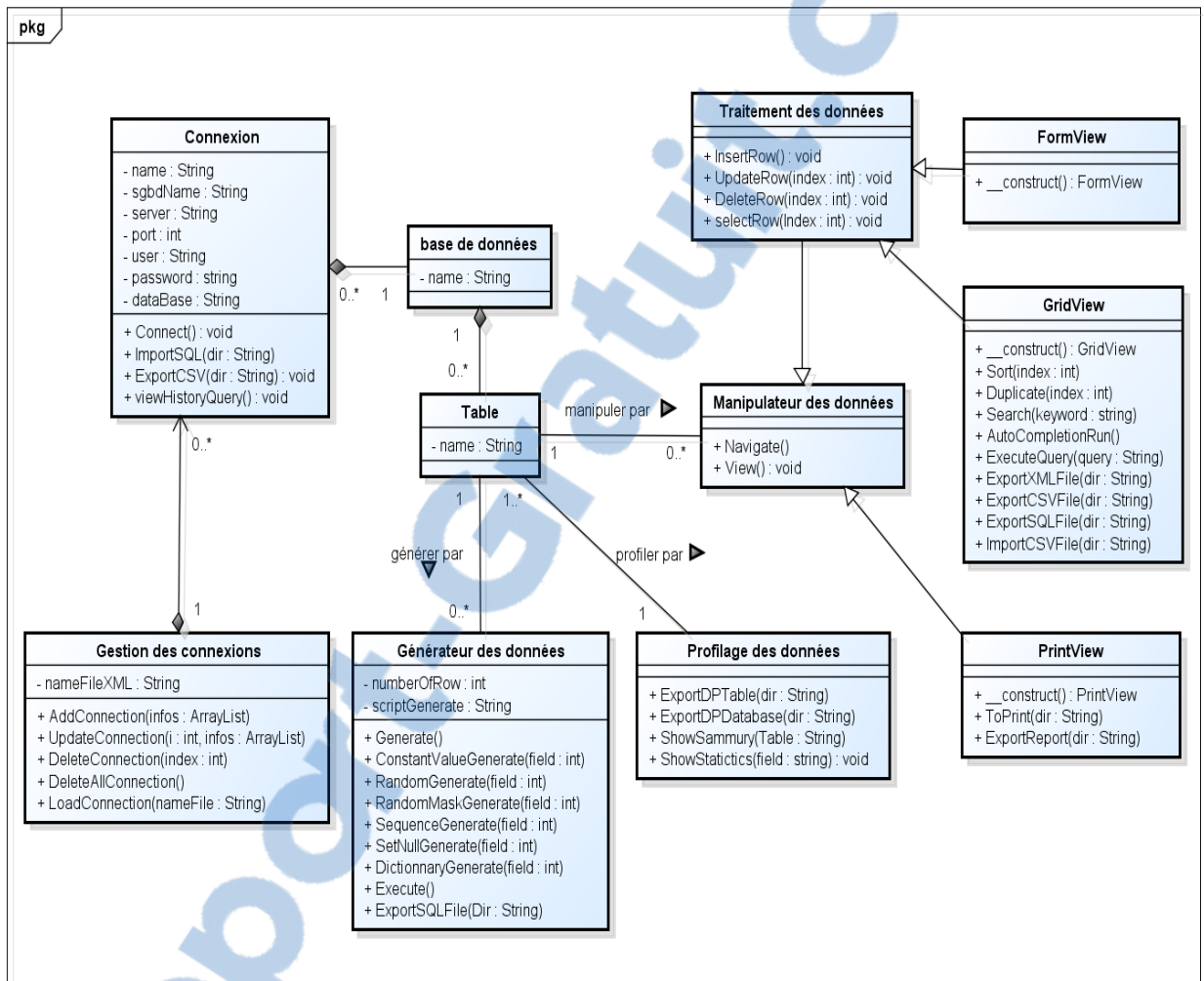


Figure II.9: Diagramme de classe.

Nous avons lié la classe connexion avec la classe gestion des connexions car cette dernière peut contenir une ou plusieurs connexions. C'est un nouveau concept que nous avons intégré dans notre application. Même chose pour la classe base de données tant que la connexion peut contenir une ou plusieurs bases de données et la base de données elle-même doit contenir à son tour une ou plusieurs tables. Cette table peut être associée à un ou plusieurs générateurs de données, profilage des données et manipulateur de données. Les tables « Print View » et « Grid View » héritent de la table « traitement de données ».

VI. Conception de l'auto-complétion

L'auto-complétion fourni aux utilisateurs une aide considérable dans l'écriture de requêtes SQL. Avant d'implémenter le système d'auto-complétion nous avons conçu un automate d'états finis. Cette automate permet de traiter les requêtes SQL d'insertion, de sélection, de modification et de suppression.

L'automate pour l'instruction « insert » :

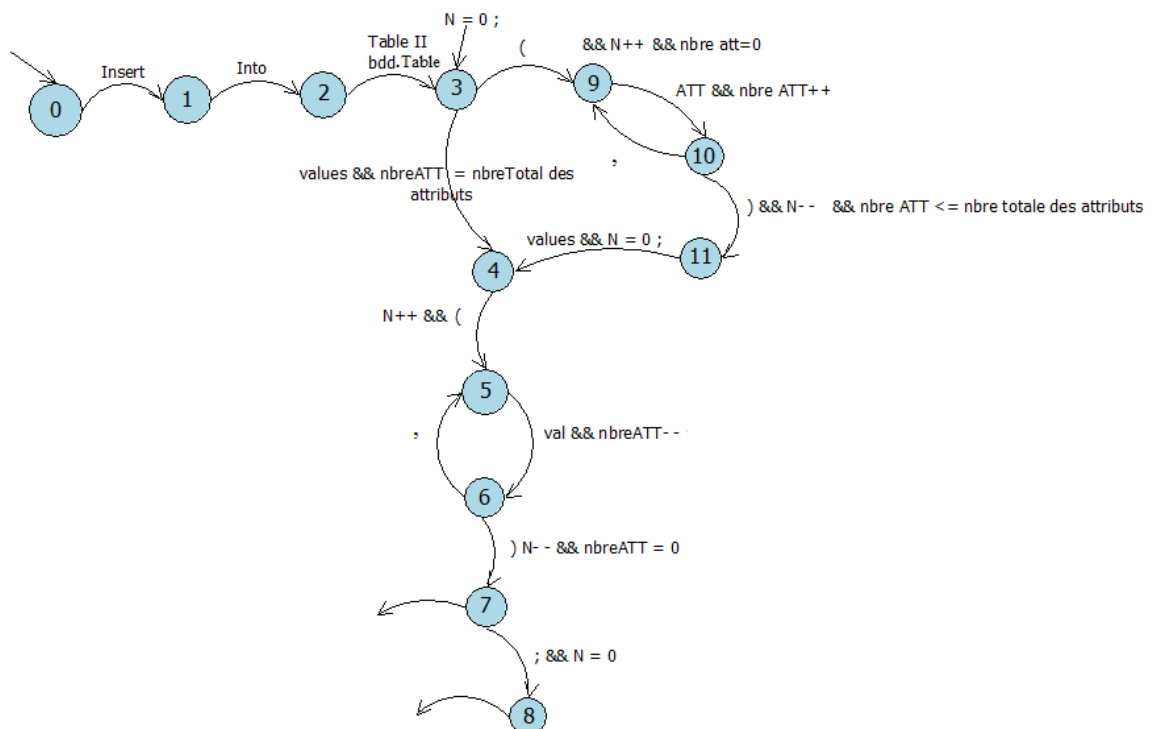


Figure II.10: L'automate pour l'instruction « insert ».

L'automate pour l'instruction « delete et update » :

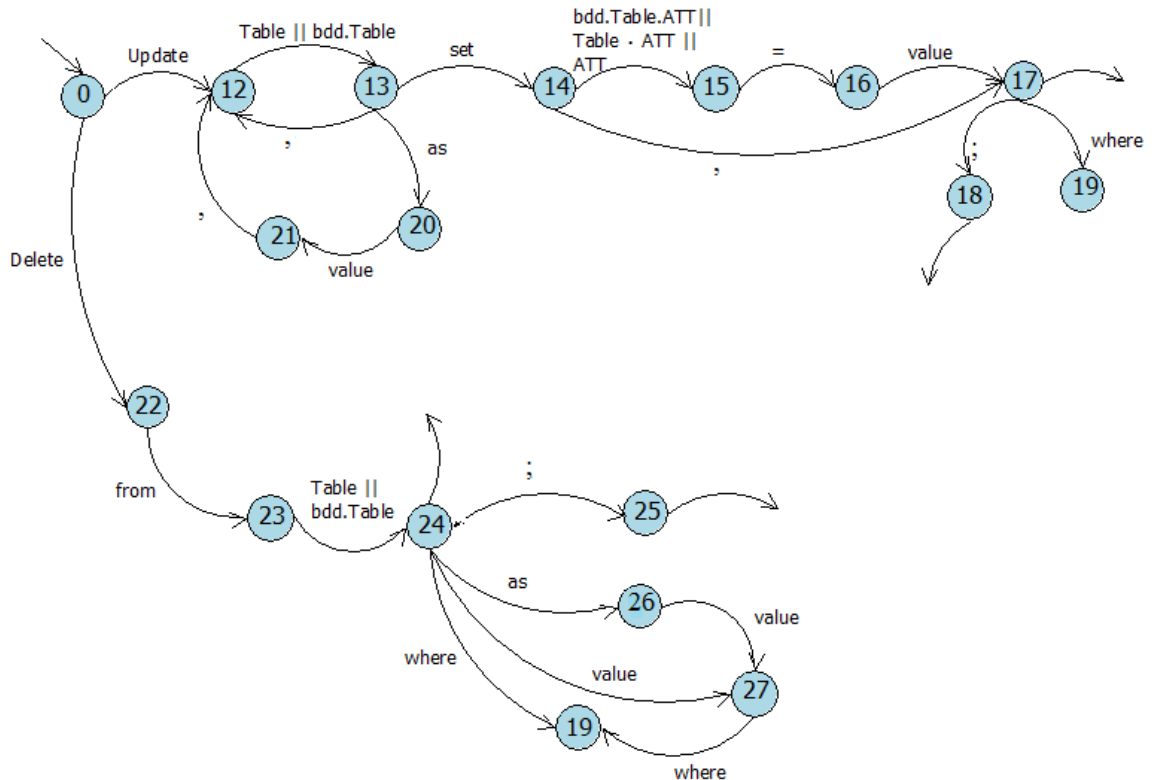


Figure II.11: L'automate pour l'instruction « delete et update ».

L'automate « Select » :

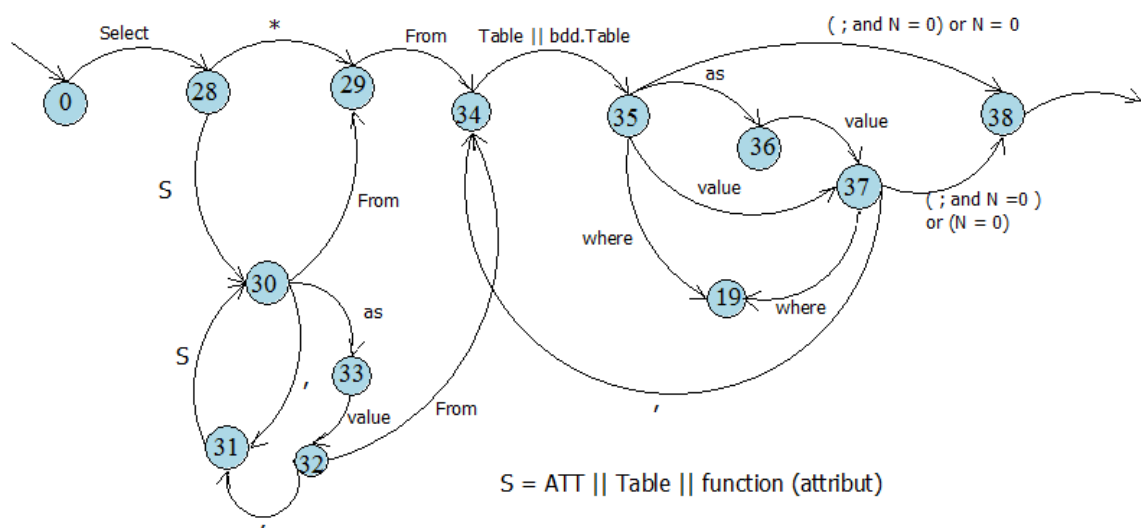


Figure II.12: L'automate pour l'instruction « select »

L'automate « Where » :

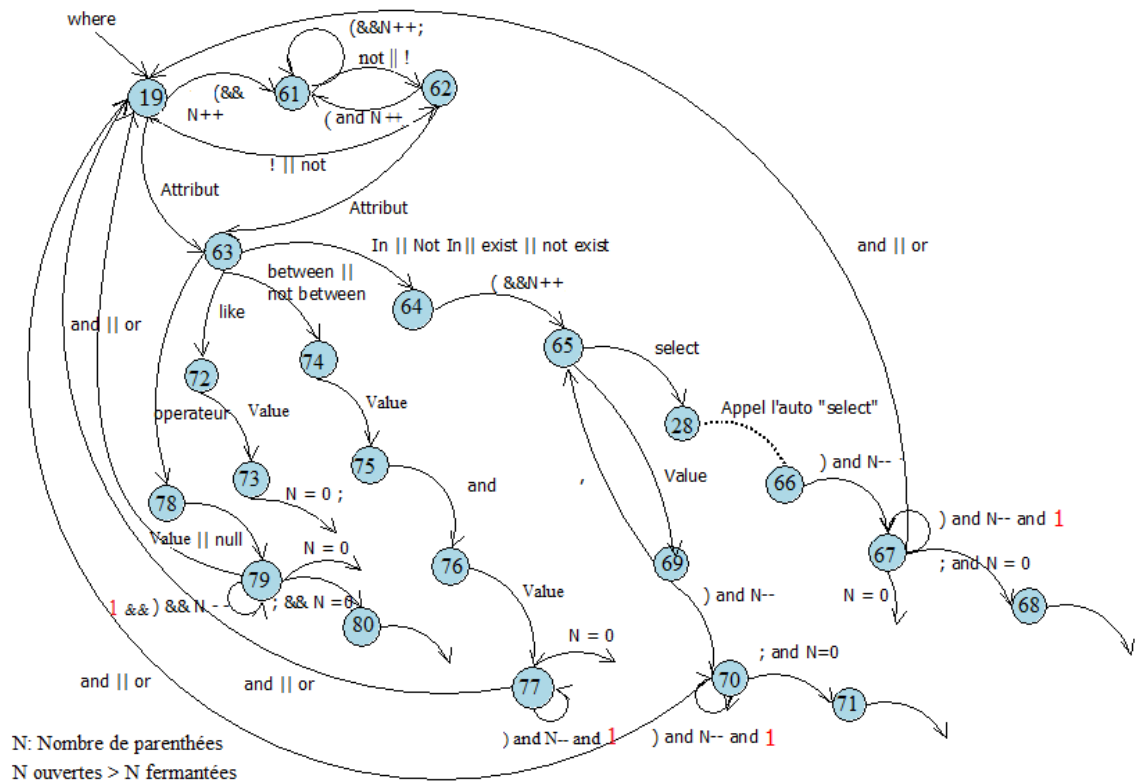


Figure II.13: L'automate pour « where ».

L'automate « Select-suite (order by, limit, group by, where) »:

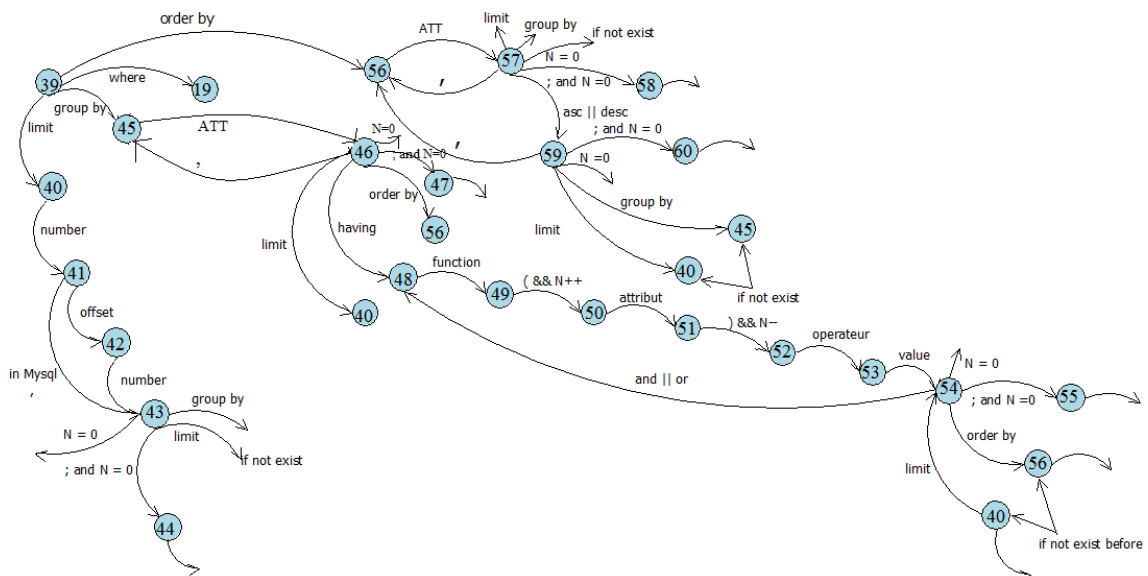


Figure II.14: L'automate pour « order by, limit, group by, where ».

VII. Conclusion

Dans ce chapitre, nous avons présenté le processus de travail que nous avons suivi tout au long de la réalisation de notre projet. Ensuite nous avons présenté l'architecture et la conception du système qui consiste à découper notre système en trois aspects : Fonctionnelle, Structurelle et Dynamique.

Les diagrammes de cas d'utilisation illustre le côté fonctionnel de notre système, ensuite. Le diagramme des classes définies la structure du système. Enfin, le côté dynamique est modélisé par les diagrammes de séquence.

CHAPITRE III

IMPLEMENTATION

I. Introduction

Après l'étape de conception de notre application, nous allons exposer dans ce chapitre la phase de réalisation, qui représente la dernière étape de cette étude.

Pour la réalisation et la mise en place de la solution, l'usage de certains outils est nécessaire pour avoir un résultat d'une application riche et fiable, comme c'est important de choisir une méthode de travail. Les outils utilisés sont présentés dans la dernière partie de ce chapitre.

II. Réalisation

Fenêtre principal :

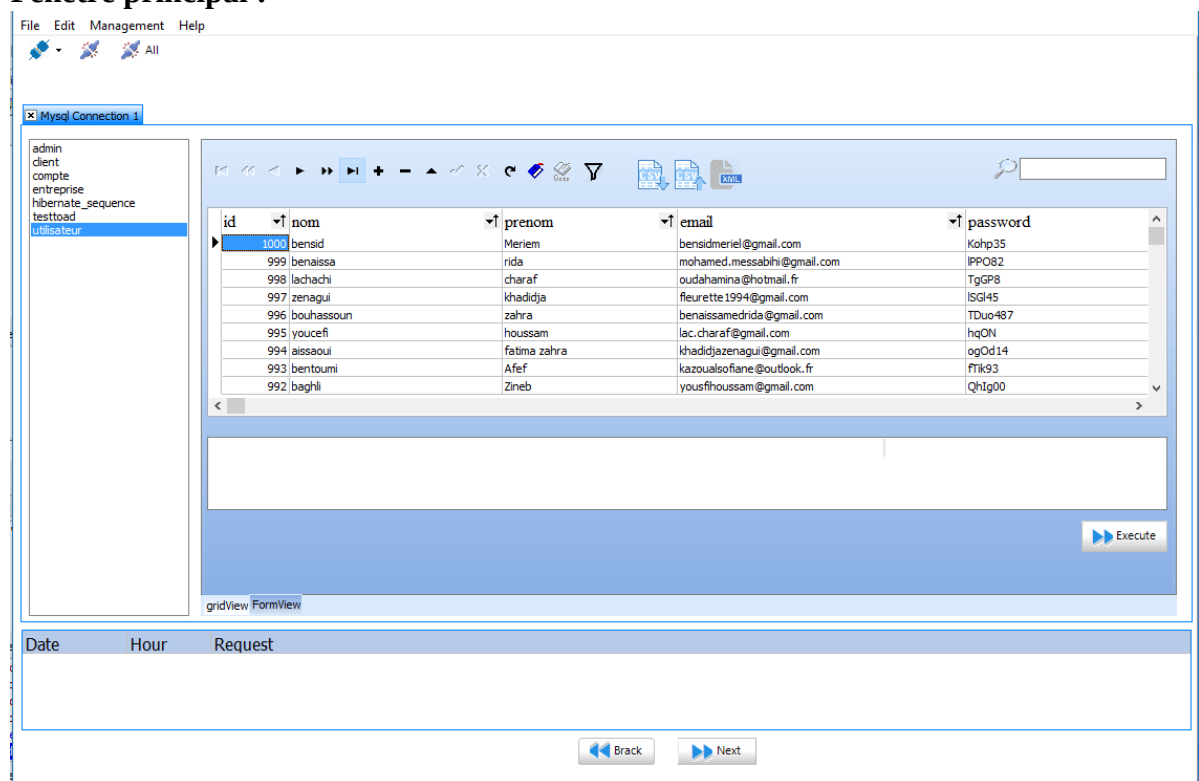


Figure III.1: Fenêtre Principale.

Dans cette figure III.1, nous allons montrer notre fenêtre principale qui contient la connexion et la gestion des connexions qui vont être bien détaillé dans les étapes suivantes plus les tables de notre base de données qui sont à gauche de la fenêtre et les données de la table sélectionnés qui sont au milieu de la page. Au-dessous de la fenêtre il y a l'historique enregistré après l'exécution d'une requête via le bouton « execute ». Comme il peut y avoir un nombre assez important des requêtes exécutés, l'utilisateur pourra cliquer sur « back » ou « next » pour recharger les requêtes SQL de l'historique.

Connexion :

Cette étape permet à l'utilisateur de remplir le formulaire ci-dessous pour ajouter une nouvelle connexion et se connecter. Au même temps les informations seront insérées dans un fichier XML si les champs sont bien remplis et nous allons avoir « successfully added » sinon la connexion ne sera pas insérée dans le fichier XML des connexions. La DTD respecté dans ce fichier XML est présenté dans la figure III.2

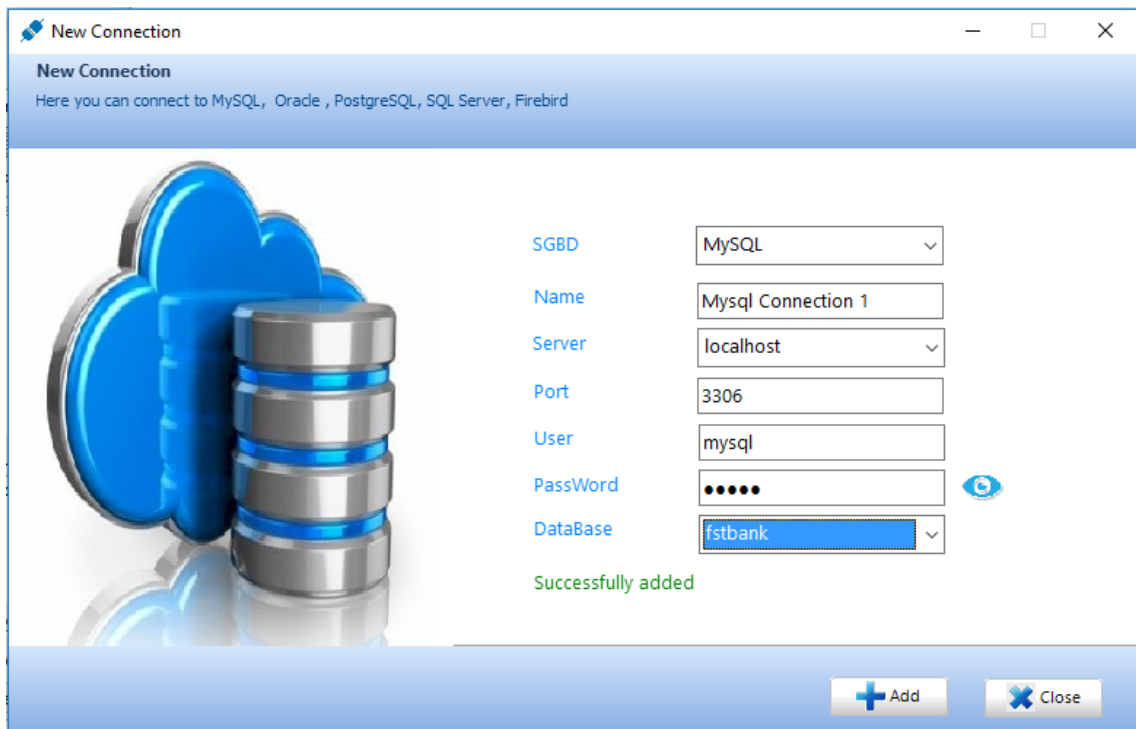


Figure III.2: Fenêtre de connexion.

Fichier DTD correspondent au fichier XML des connexions :

```
<!ELEMENT Connections(Connection*)>
<!ELEMENT Connection(name, sgbd, server, port, user Name, pass, dataBase)>
<!ELEMENT name (#PCDATA)>
<!ELEMENT sgbd (#PCDATA)>
<!ELEMENT server (#PCDATA)>
<!ELEMENT port (#PCDATA)>
<!ELEMENT userName (#PCDATA)>
<!ELEMENT pass (#PCDATA)>
```

Figure III.3: DTD de fichier XML des connexions.

Gestion des connexions :

La liste des connexions est affichée dans la figure III.4. Ci-dessous un ensemble de fonctionnalités qui aide à la gestion de nos connexions :

- L'ajout d'une nouvelle connexion par l'utilisateur lui-même et nous avons déjà parlé de cette dernière dans l'étape précédente ;
- La suppression d'une connexion quelconque selon le besoin de l'utilisateur soit une seule est sélectionnés par lui-même sinon il peut effectuer un « delete-all » pour éliminer tous les connexions existantes ;
- Comme il peut faire la modification en cas où il se trompe ou il veut changer le nom de sa connexion ;

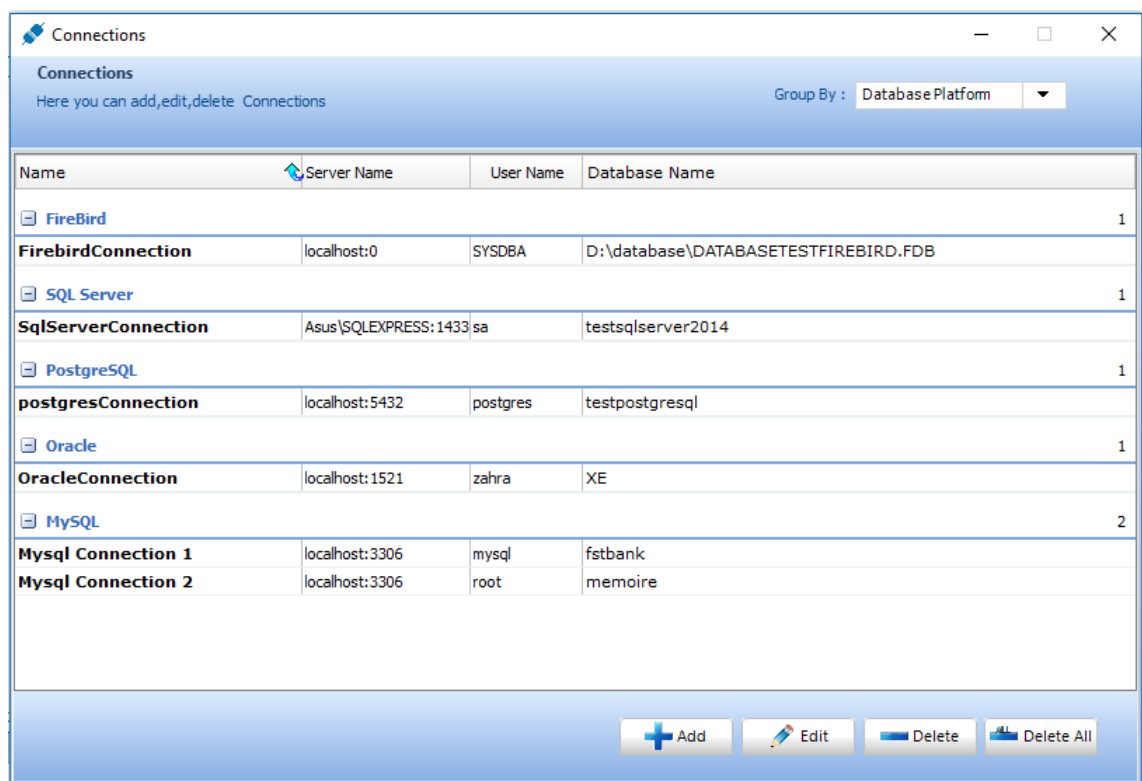


Figure III.4: Fenêtre de gestion des connexions.

Exécution de requêtes en multithreads :

Notre système permet d'ouvrir plusieurs connexions à plusieurs bases de données en parallèle. Chaque requête qui est exécutée dans une connexion est traitée par un thread à part. Ceci permettra d'exécuter plusieurs requêtes SQL en parallèle. La figure ci-dessous montre cinq connexions ouvertes en parallèle.

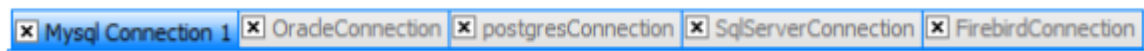


Figure III.5: Capture de multi connexions.

GridView :

Cette fenêtre est très riche en point de vue fonctionnalités. Après que l'utilisateur s'authentifie, notre système affichera cette dernière. Puis en cliquant sur la table, les données seront affichées comme il est montré dans la figure « » pour qu'il puisse naviguer dans sa table en modifiant, supprimant ou ajoutant des données. Par la suite il peut trier ses données (ascendant/descendant). En plus il pourra exécuter n'importe quelle requête (select, delete, update, insert) avec l'aide de l'auto-complétion que nous avons implémenté. À droite il peut faire une recherche (filtrage des données), il y a aussi l'actualisation de la page et l'historique des requêtes. En plus il peut faire tout type d'export/import (CSV,SQL, XML).

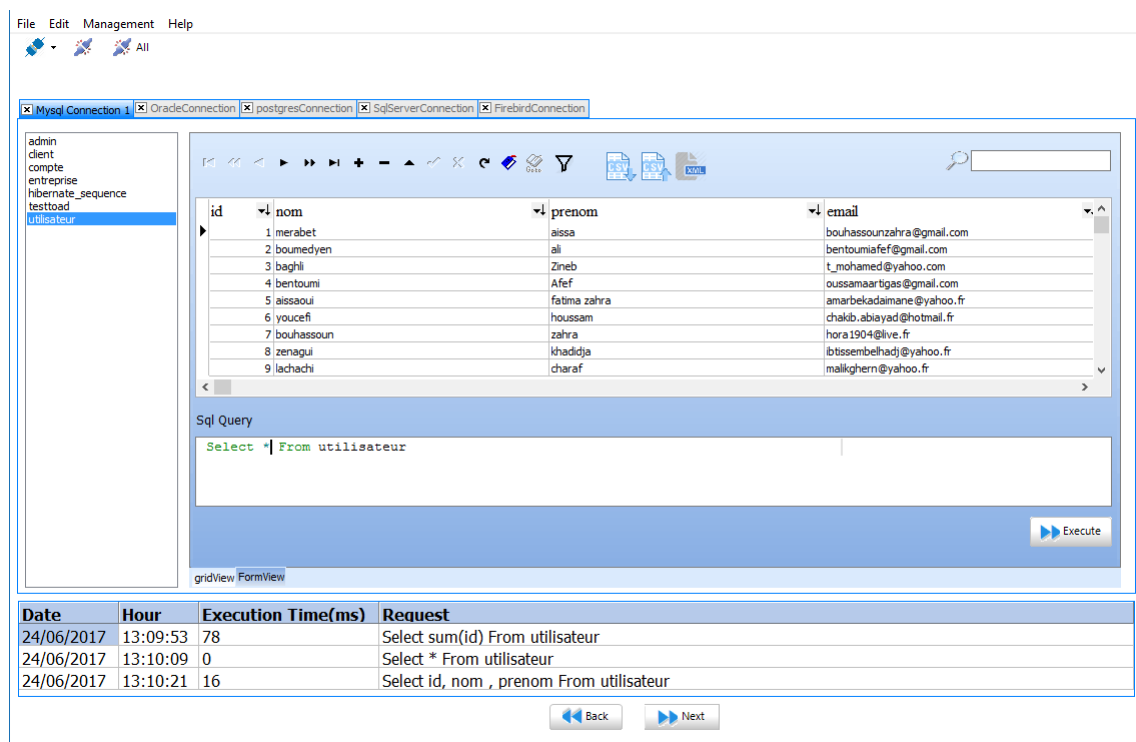


Figure III.6: Fenêtre Grid View.

Exportation en XML :

Une fois que l'utilisateur saisie une requête et l'exécute, il peut par la suite exporter les informations présentes dans le GridView dans un fichier XML. Ci-dessous un exemple de ce fichier.

[illegible]

Figure III.7: Fichier d'exportation en XML.

Voici le DTD correspondant au fichier xml :

```
<!--ELEMENT Database [SQLScript, Records ]>
<!--ELEMENTSQLScript(#PCDATA)>
<!--ELEMENT Records(Record*)>
<!--ELEMENT Record(field+)>
<!--ATTLIST field name CDATA #REQUIRED >
```

Figure III.8: DTD de fichier d'export en XML.

Exportation en CSV :

La figure III.9 montre l'export en CSV. L'utilisateur a le droit de faire un export total des données comme il peut aussi choisir ses champs par le biais de « add » et même chose pour la suppression des champs soit il effectue une illimitation totale ou partielle. Nous avons offert de plus le choix du séparateur comme il s'est illustré dans la figure ci-dessous en bas de la fenêtre et juste à côté, à droite il a deux possibilités soit il enregistre en cliquant sur « save » l'export, soit il annule en cliquant sur « cancel ».

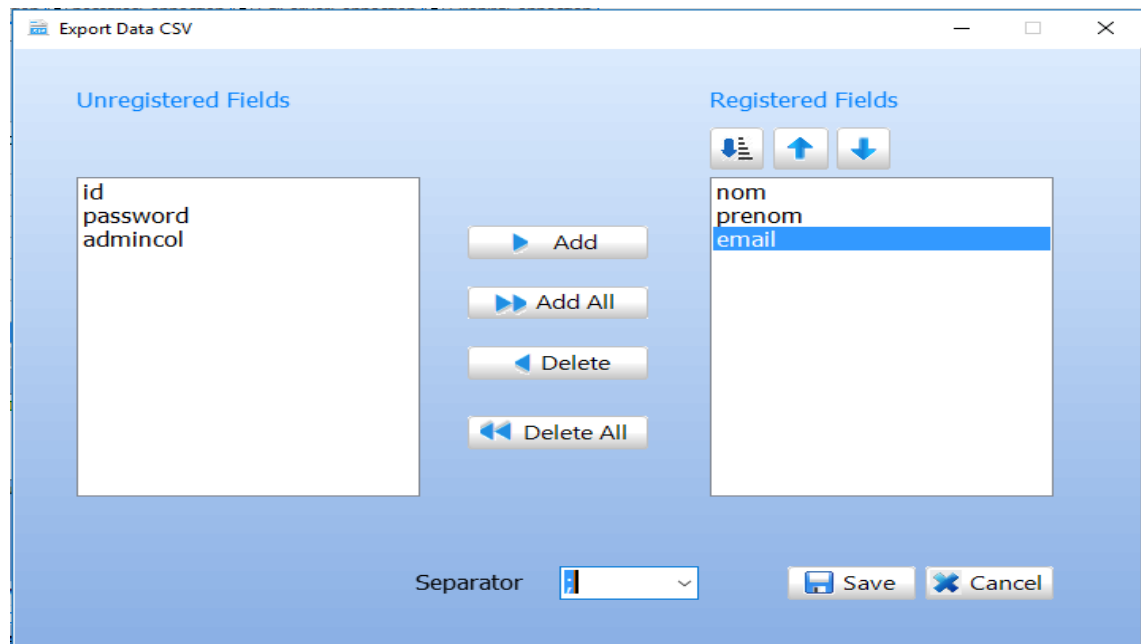


Figure III.9: Fichier d'exportation en CSV.

Exemple d'un export en CSV :

	A	B	C
	prenom	nom	email
1	zahra	bouhassoun	bouhassounzahra@gmail.com
2	afef	bentoumi	bentoumiafef@gmail.com
3	mohamed	tadlaoui	t_mohamed@gmail.com
4	hanane	bouhassoun	b_hanane2000@gmail.com
5	fatima zahra	aissaoui	amarbekadaimane@yahoo.fr
6	houssam	youcefi	chakib.abiayad@hotmail.fr
7	zahra	bouhassoun	hora1904@live.fr
8	Meriem	bensid	sofiane_belhadj_kacem@outlook.fr
9	houari	mahfoud	h_mahfoud@gmail.com

Figure III.10: Exemple d'exportation en CSV.

Dans cette figure ci-dessus nous allons présenter un fichier CSV après l'export des données pour montrer le résultat.

Form View :

Cette fenêtre a les mêmes caractères de la fenêtre « GridView » sauf au lieu que le système affichera la liste des données, il affichera ici les champs de l'enregistrement sélectionné de la table désigné. L'intérêt ici c'est d'insérer/supprimer une données en cliquant sur valider afin de les valider. Comme nous avons enlevé certaines fonctionnalités par exemple : la recherche, l'export, le tri et l'exécution des requêtes etc.

Date	Hour	Execution Time(ms)	Request
24/06/2017	13:09:53	78	Select sum(id) From utilisateur
24/06/2017	13:10:09	0	Select * From utilisateur
24/06/2017	13:10:21	16	Select id, nom , prenom From utilisateur

Figure III.11: Fenêtre Form View.

PrintView:

Dans cette figure, nous avons montré un exemple de « PrintView » c.-à-d. dans le cas où l'utilisateur veut imprimer ses données il fait un clic droit sur sa table qui se trouve dans la page « PrintView » et en choisissant par la suite « print » pour avoir la figure III.12 ci-dessous.

id	nom	prenom	email
1	bouhassoun	zahra	bouhassounzahra@gmail.com
2	bentoumi	afef	bentoumiafef@gmail.com
3	tadlaoui	mohamed	t_mohamed@gmail.com
4	bouhassoun	hanane	b_hanane2000@gmail.com
5	aissaoui	fatima zahra	amarbekadaimane@yahoo.fr
6	youcefi	houssam	chakib.ablayad@hotmail.fr
7	bouhassoun	zahra	hora1904@live.fr
8	bensid	Meriem	sofiane_belhadj_lacem@outlook.fr
9	mahfoud	houari	h_mahfoud@gmail.com
10	benalissa	reda	samistar2a@gmail.com
11	bensid	Meriem	sofiane_belhadj_lacem@outlook.fr
12	ouadah	Amina	merabetalissa1982@gmail.com
13	sekkak	Oussama	boumedyen121@gmail.com
14	tadlaoui	mohamed	bag_zineb@hotmail.fr
15	hamdaoui	sara	bentoumiafef@gmail.com
16	benamar	adam	ais-fatima-z@outlook.fr
17	mscsahhi	ismail	youcfihoussam@gmail.com

Figure III.12: Fenêtre Print View.

Generate Data:

Cette fonction offre la génération des données qui est une étape très importante car chaque développeur aura besoin de tester son travail en peuplant la table automatiquement sans aucun effort et avec des milliers ou des millions d'enregistrement de différents types et cas. L'utilisateur est censé de sélectionner en premier lieu le nombre de ligne à générer comme il est montré en haut de la figure III.13, plus qu'il doit effacer les données qui existe en cochant sur « clear existing data ». Après il effectue le choix des types pour chaque champ par exemple le champ « prénom » est de type var-char, et comme action en utilisant le dictionnaire il choisit « First Name ». Par la suite il a deux possibilités soit il effectue un export de ses données sous format SQL. Sinon il exécute directement cette opération en cochant « execute now » qui est montré dans la figure III.13 en bas de la page « Generate Data ». Après il clique sur « Generate » qui se trouve dans la figure III.13, À droite de la page. Comme il peut annuler la génération.

The screenshot shows a 'Generate Data' window with three steps:

- Step 1:** Number of rows to generate: 100. ☒ Clear existing Data.
- Step 2:** Options for column 'id'. Column Type: integer. Action: Sequence generate. Start: 1. Step: 1. A 'Reset' button is present.
- Step 3:** ☐ Export Script in file sql. ☒ Execute now.

Buttons: Generate, Cancel.

Figure III.13: Fenêtre la génération des données.

Profiling des données

Dans cette page, nous avons intégré trois types d'analyses (structurelle, de BDD, de colonne). Le système offre le « Data Profiling » c.-à-d. les statistiques d'une base de données. Sur l'onglet « Statistics » qui est présent à gauche et en haut de la figure III.14 le système affiche des statistiques de chaque champ de la table d'une base de données (data-type, min, max, sum....), aussi comme il est montré dans la figure suivante :

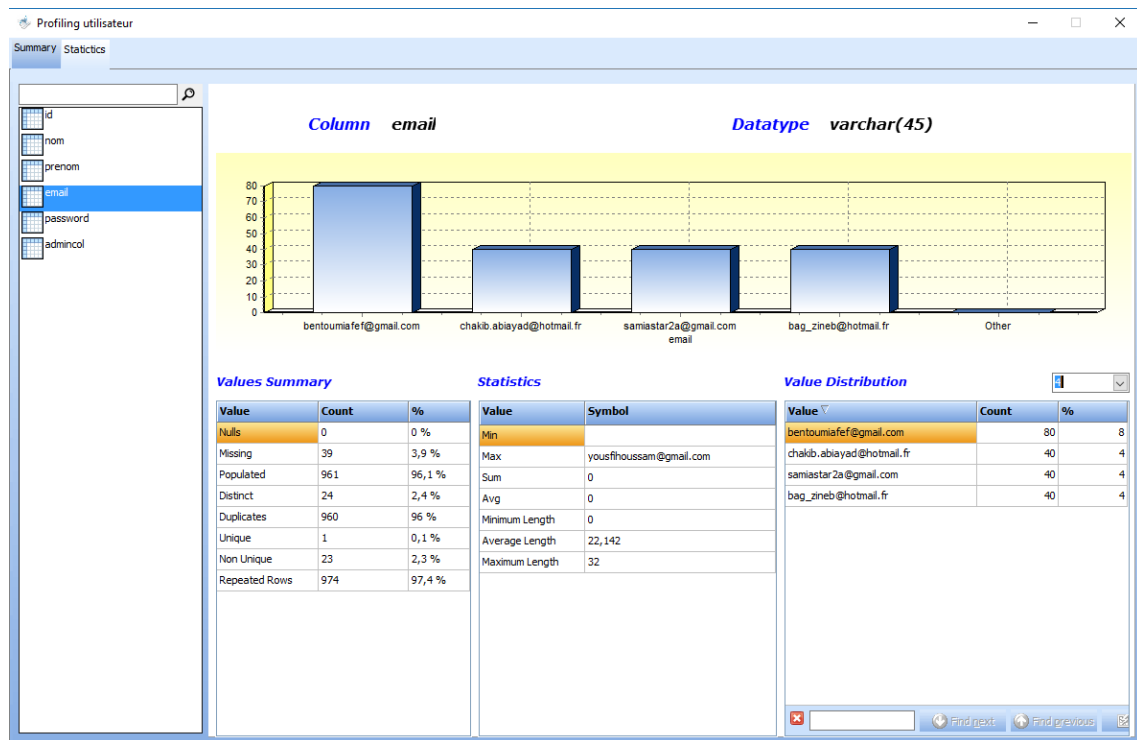


Figure III.14: Fenêtre data profiling « Statistic ».

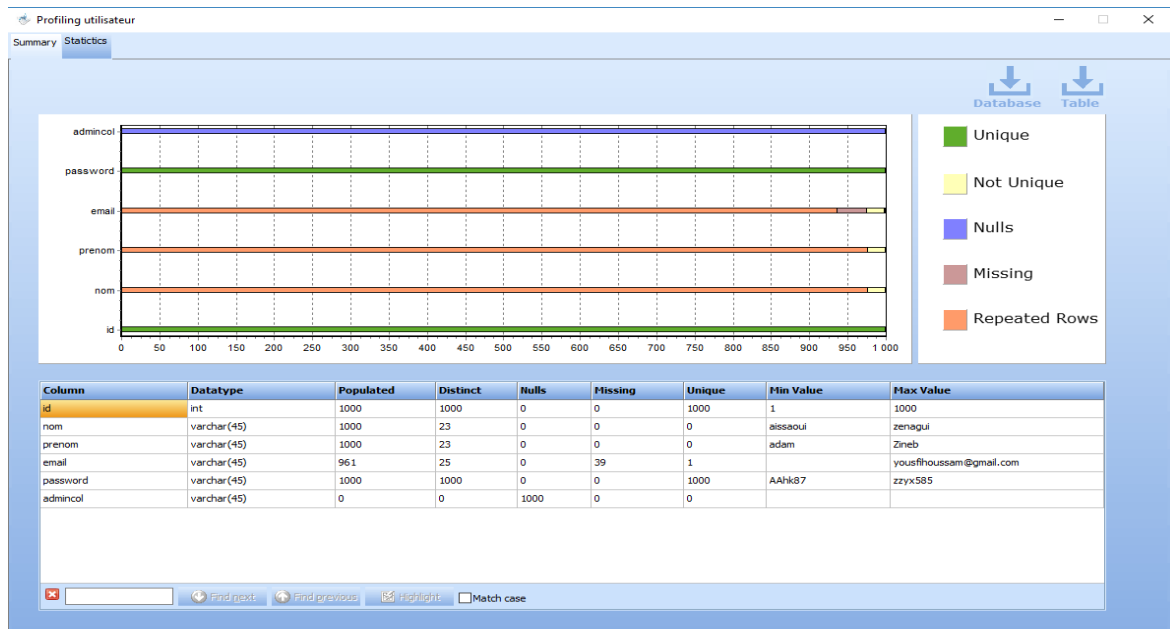


Figure III.15: Fenêtre data profiling « Statistic ».

Et voilà la figure III.15 ci-dessus qui présente l'onglet « summary » qui affiche les statistiques (Unique, Null, Distinct, Repeated Rows etc.). Ici l'utilisateur pourra télécharger une version des statistiques de sa table ou totalement de sa base de données, comme il est montré en haut de la page à droite.

Rapports de Data Profiling :

Cette figure III.16 montre un rapport qui récapitule les statistiques générées d'une base de données. Ce rapport web dynamique permet à l'utilisateur de naviguer dans les différentes statistiques de la base de données.

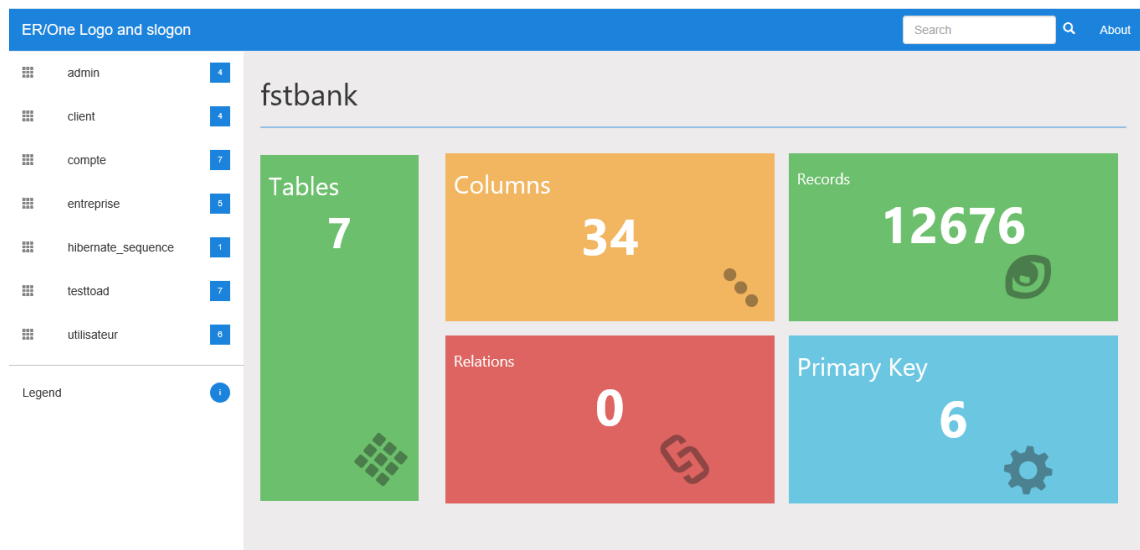


Figure III.16: Rapport d'une BDD

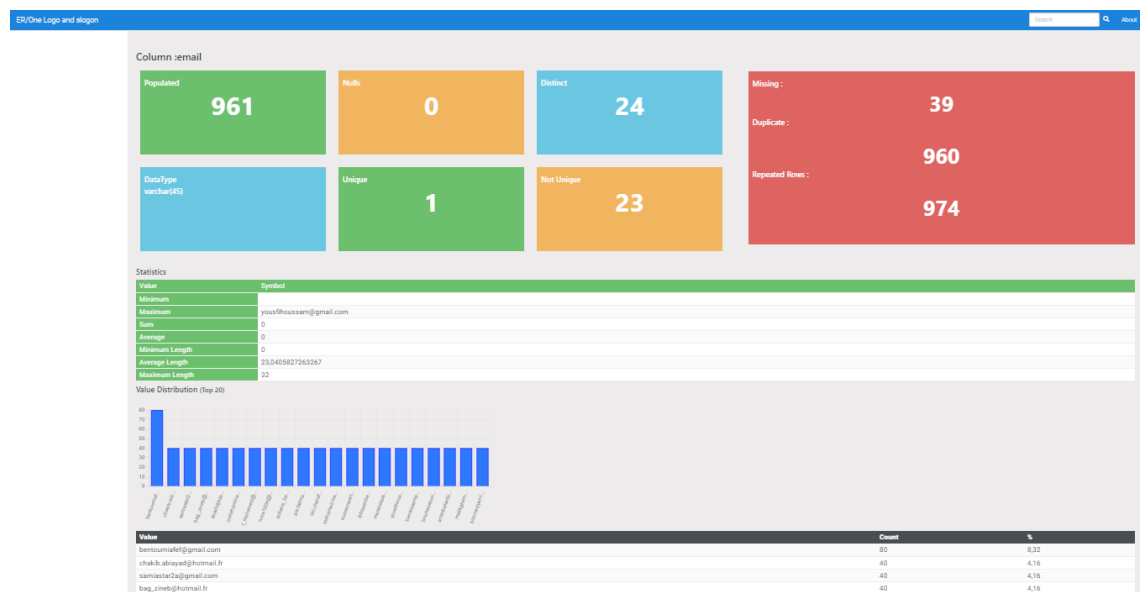


Figure III.17: Rapport d'un champ

III. Gestion des tests et anomalies

1. Tests

Les tests sont nécessaires pour détecter le plus rapidement possible les problèmes logiciels ou les bugs de manière à pouvoir les corriger avant la publication et le déploiement de l'application. Ils sont également réalisés pour développer la confiance et la connaissance des applications. Nous utilisons ces tests, pour valider l'application, vérifié si tout fonctionne comme prévu et si elle répond aux exigences.

2. Anomalies

Une application informatique a forcément quelques erreurs. Ces anomalies sont mises en évidence par la phase des tests cités ci-dessus. Il est important d'associer à chaque déclaration d'anomalie un type permettant de mesurer la qualité globale de notre application. Pour cela nous avons classé les anomalies en trois types : bloquante, majeure et mineure. Nous avons préconisé de saisir les bugs d'une façon ressemblante à un scénario pour que le correcteur (développeur) puisse les comprendre, les reproduire et les corriger.

- **Anomalie bloquante** : anomalie rendant impossible l'utilisation d'une ou plusieurs fonctionnalités du système. Il s'agit d'erreurs graves pouvant bloquer tout le système. Exemple : Génération de données une fois le nombre des données dépasse les 100000 lignes.

- **Anomalie majeure**: toute anomalie autre que bloquante qui altère le fonctionnement de tout ou de l'une des fonctionnalités principales. Exemple lorsqu'on clic sur le bouton « execute » sur l'éditeur de texte SQL il reste vide ou il affiche « requête invalide » implique l'affichage d'un message (rien à exécuter ou syntaxe incorrect).

- **Anomalie mineure** : elle diffère des deux types cités ci-dessus. Il s'agit d'erreurs qui peuvent être gênantes, mais ne présentent pas de danger pour le système.

Exemple :

- Le bouton échappe permet de créer des tab dans la page contrôle du menu principale (il faut renommer les tab par nom de la base de données (avec ou sans le type du SGBD)).
- dans la partie summary (data profiling) si nous cherchons un champ et cliquons sur highlight alors nous avons tombé sur ce bug le nom de la colonne est dupliqué dans les statistiques.

IV. Outils de développement

IV.1 Gestion de projet « Trello »

Trello est un outil de gestion de projet en ligne et ergonomique. Nous l'utilisons pour planifier et organiser nos tâches entre différents membres de notre équipe. Grâce à cet outil nous simplifions le suivi de notre projet.

Trello se présente sous la forme d'un tableau de bord (boards) et chaque tableau comporte des listes de cartes. Trello est dessiné pour répartir des listes de tâches sous forme de colonnes « To Do », « Doing » et « Done ». Ces tâches sont assignées aux différents membres du projet.

Au fil de leur exécution, il suffit de glisser-déposer les cartes correspondantes d'une colonne à l'autre. Des codes couleurs permettent de gérer les priorités. Trello propose aussi des dates limites et des notifications pour ne manquer aucune étape. Chaque sprint de notre projet est associé à un tableau Trello, Ci-dessous un exemple du tableau du premier sprint.

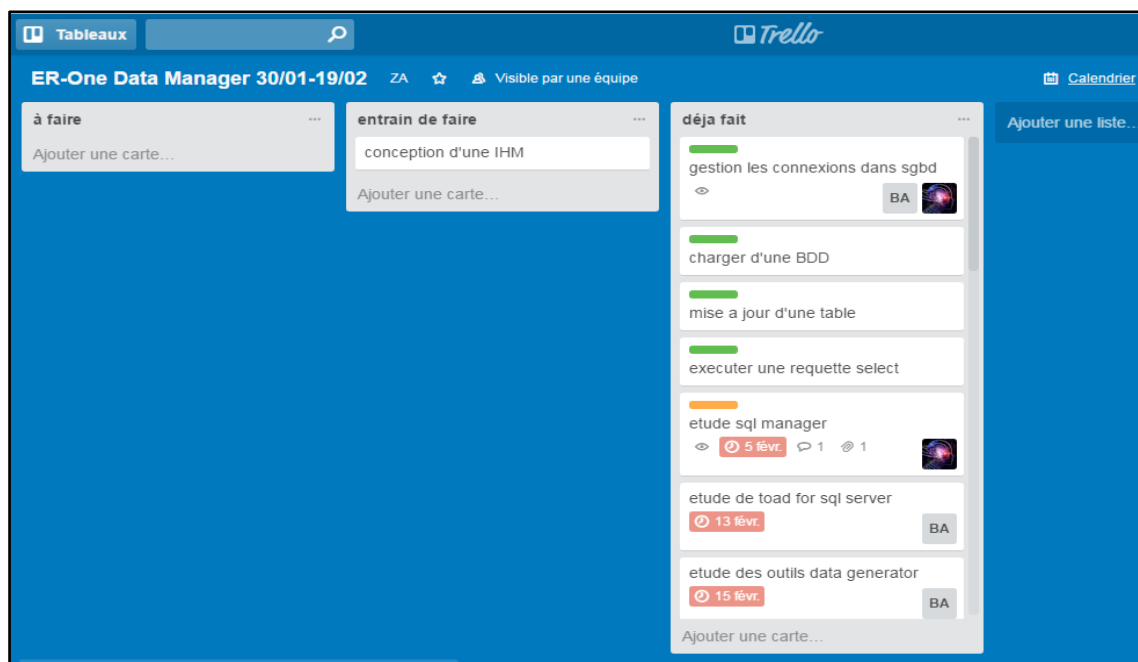


Figure III.18: Une vue sur Trello.

IV.2 Modélisation « astah »

Astah est un outil de modélisation UML créé par la compagnie japonaise ChangeVision, il est gratuit, d'une simple utilisation, intuitif pour un débutant. Astah supporte la norme UML.

Cet outil propose les diagrammes UML nécessaires à une bonne modélisation.

IV.3 Programmation « Embarcadero Delphi »

Embarcadero Delphi est un kit de développement logiciel (SDK) pour applications bureautiques, mobiles, web et console. Les compilateurs de Delphi utilisent leur propre dialecte Pascal d'objet et génèrent un code natif pour plusieurs plates-formes: Windows (x86 et x64), OS X (32 bits uniquement), iOS (32 et 64 bits), Android et Linux (64 bits Intel).

Delphi, une partie de RAD Studio, comprend un éditeur de code avec Code Insight (exécution du code), Error Insight (vérification en temps réel des erreurs) et d'autres fonctionnalités; Factorisation; Un concepteur de formulaires visuels pour VCL (native Windows) et FMX (plate-forme croisée, partiellement native par plate-forme); Un débogueur intégré pour toutes les plates-formes, y compris le mobile; Contrôle source (SVN, git et Mercurial); Et un support pour les plug-ins tiers. Il possède un solide support de base de données. Il n'est pas inhabituel qu'un projet Delphi d'un million de lignes compile en quelques secondes - un benchmark a donné 170 000 lignes par seconde. Il est en développement actif, avec des versions (en 2016) tous les six mois, avec de nouvelles plates-formes ajoutées approximativement à chaque seconde version.

Delphi a été initialement développé par Borland comme un outil de développement rapide d'applications pour Windows en tant que successeur de Turbo Pascal. Delphi a ajouté une orientation complète de l'objet à la langue existante, et depuis, la langue a augmenté et prend en charge de nombreuses autres fonctionnalités de langage moderne, y compris les méthodes génériques et anonymes, ainsi que des fonctionnalités inhabituelles telles que des types de chaînes intégrées et un support COM natif. Delphi et son équivalent C ++, C ++ Builder, partagent de nombreux composants principaux, notamment l'IDE, la Visual Component Library (VCL) et une grande partie de la RTL, et sont compatibles les uns avec les autres: C ++ Builder 6 et versions ultérieures peuvent consommer Les fichiers de langue Delphi et C ++ dans le même projet et les paquets compilés avec C ++ Builder écrit en C ++ peuvent être

utilisés à partir de Delphi. En 2007, les produits ont été diffusés conjointement comme RAD Studio. RAD Studio est un hôte partagé pour Delphi et C ++ Builder, et peut être acheté avec l'un ou l'autre [20].

V. Programmation

Pour la réalisation de notre application, nous utilisons les langages de programmation suivants : Delphi, SQL, XML, HTML, CSS, BootStrap, JavaScript. et comme environnement de travail, nous avons utilisé Delphi intégré (EDI) développé par Borland. Ce dernier nous a permis de développer d'une manière plus efficace et propre.

VI. Conclusion

La réalisation d'une application nécessite une maîtrise du domaine et une maîtrise des outils pour mettre en place un environnement d'exécution fiable. Dans notre travail nous avons utilisé plusieurs langages tels-que delphi, xml, html et javascript. En plus il faut un bon esprit d'équipe et des compétences dans la gestion des projets pour gérer la distribution des tâches et résoudre les conflits. Pour gérer notre projet nous avons utilisé la méthode Scrum et l'outil Trello. Ces différents langages, méthode et outils nous ont permis de réaliser trois modules applicatifs qui permettent l'administration et la manipulation des données relationnelles.

CONCLUSION GENERAL

Dans ce projet, nous avons eu la chance de réaliser trois outils. Ces derniers manipulent les cinq SGBDs (Oracle, MySQL, SQL Server, PostgreSQL, Firebird). Ils servent à la manipulation des données qui désigne toute chose concernant le traitement des données, leur navigation et même leur présentation. Ensuite nous avons créé aussi l'outil de génération de données pour le remplissage des données afin de les utiliser dans les tests des applications logicielles. En plus, nous avons réalisé un outil de data profiling qui offre des statistiques et des informations sur la table et la BDD. Nous avons pu aussi aboutir les objectifs fixés dans le cahier de charges c.-à-d que nous avons validé la solution et que nous avons satisfait les besoins de la société Soft Builder.

Notre stage de fin d'étude nous a donné l'occasion de mieux comprendre le domaine des BDD en fournissant notre propre solution. Nous avons eu la chance aussi d'appliquer nos connaissances théoriques sur le terrain, voir le résultat de notre effort, tout en recevant un retour du côté d'entreprise et les clients pour améliorer notre solution. Donc c'est un processus continu d'apprentissage, d'amélioration et d'innovation.

Nous avons compris que concevoir une solution et développer une application est un travail d'équipe et il nécessite plusieurs compétences en termes de gestion, de planification, de communication et bien sur des compétences de développement.

Le travail dans ce projet est loin d'être finis, il reste toujours un espace pour l'amélioration et l'ajout d'autres fonctionnalités, ce qui est prévu pour le prochain sprint. Comme perspective nous comptons terminer l'intégration d'auto-complétion, intégrer d'autres SGBDs, avoir des statistiques plus poussées dans le data profiling et améliorer l'IHM.

REFERENCES BIBLIOGRAPHIQUES

[1] :http://www.memoireonline.com/10/13/7468/m_Linformatisation-de-la-gestion-des-abonnes-de-la-SNEL--Societe-nationale-deelectricite-en4.html , consulté le 01/05/2017.

[2] :http://www.ird.fr/informatique-scientifique/documents/sghd/sql/formation_SGBD_cours_SGBD.pdf, consulté le 01/05/2017.

[3] :<https://fr.scribd.com/document/73272955/Decouverte-MySQL-PostgresSQL-Oracle>, consulté le 04/05/2017.

[4] :https://fr.wikipedia.org/wiki/Oracle_Database , consulté le 04/05/2017.

[5] :http://www.ird.fr/informatique-scientifique/documents/sghd/sql/formation_SGBD_cours_SGBD.pdf , consulté le 05/05/2017.

[6] :SQL Server 2008: Administration d'une base de données avec SQL Server Management Studio, Jérôme Gabillaud Février 2009.

[7] :Utiliser des bases de données de Firebird dans le système d'informations des entreprises Philippe Makowski 15 Juillet 2006.

[8] :https://fr.wikipedia.org/wiki/MySQL_Workbench, consulté le 09/02/2016

[9] :<http://www.sqlmanager.net/fr/products/mysql/manager> , consulté le 09/02/2016.

[10] :[https://fr.wikipedia.org/wiki/Toad_\(logiciel\)](https://fr.wikipedia.org/wiki/Toad_(logiciel)) , consulté le 10/02/2017.

[11] :<https://www.quest.com/fr-fr/products/toad-for-sql-server/>, consulté le 10/02/2017.

[12] :<https://help.talend.com/reader/W5EvAVIgByIFRYjKp6Di9g/MDauhq9nl1m0FXZX~aHp9Q> , consulté le 19/02/2016

[13] <http://www.sqledit.com/dg/> , consulté le 22/02/2016.

- [14] :<https://www.techopedia.com/definition/30405/test-data-generator> , consulté le 24/02/2017.
- [15] :<http://featurist.com.au/feature-it-mockaroo/> , consulté le 03/03/2017.
- [16] :Claude Aubry, SCRUM, le guide pratique de la méthode agile la plus populaire, Préface de François Beauregard, Dunod, 2011. Consulté le 01/03/2017.
- [17] :BELHABIB A., MATAHRI A., Conception d'un système de recommandation pour un réseau sociale d'apprentissage, projet de fin d'étude, université de Tlemcen, 2014, consulté le 01/03/2017.
- [18] :Claude Aubry, SCRUM, le guide pratique de la méthode agile la plus populaire, Préface de François Beauregard, Dunod, 2011. Consulté le 01/03/2017.
- [19] :<https://openclassrooms.com/courses/debutez-l-analyse-logicielle-avec-uml/uml-c-est-quoi> , consulté le 05/06/2017.
- [20] :[https://en.wikipedia.org/wiki/Delphi_\(programming_language\)](https://en.wikipedia.org/wiki/Delphi_(programming_language)), 15/06/2016.

LISTE DE FIGURES

Figure II.1: Burn Down Chart: Iteration.....	25
Figure II.2: Architecture.....	28
Figure II.3: Diagramme de cas d'utilisation.....	29
Figure II.4: Diagramme de séquence « Connexion ».....	30
Figure II.5: Diagramme de séquence «gestion des Connexions».....	31
Figure II.6: Diagramme de séquence «Manipulation des données».....	33
Figure II.7: Diagramme de séquence «Génération des données».....	34
Figure II.8: Diagramme de séquence «Data Profiling».....	35
Figure II.9: Diagramme de classe.....	36
Figure II.10: L'automate pour l'instruction « insert».....	37
Figure II.11: L'automate pour l'instruction « delete et update».....	38
Figure II.12: L'automate pour l'instruction « select».....	38
Figure II.13: Automate pour « where».....	39
Figure II.14: L'automate pour « order by, limit, group by, where».....	39
Figure III.1: Fenêtre Principale.....	42
Figure III.2: Fenêtre de connexion.....	43
Figure III.3: DTD de fichier XML des connexions.....	43
Figure III.4: Fenêtre de gestion des connexions.....	44
Figure III.5: Capture de multi connexions.....	45
Figure III.6: Fenêtre Grid View.....	45
Figure III.7: Fichier d'exportation en XML.....	46
Figure III.8: DTD de fichier d'export en XML.....	46
Figure III.9: Fichier d'exportation en CSV.....	47
Figure III.10: Exemple d'exportation en CSV.....	47
Figure III.11: Fenêtre Form View.....	48
Figure III.12: Fenêtre Print View.....	49
Figure III.13: Fenêtre la génération des données.....	50
Figure III.14: Fenêtre data profiling « Statistic ».....	51
Figure III.15: Fenêtre data profiling « Statistic ».....	51
Figure III.16: Rapport d'une BDD.....	52
Figure III.17: Rapport.....	52
Figure III.18: Une vue sur Trello.....	54

LISTE DES TABLEAUX

Tableau I.1 Tableau comparatif de manipulation des données.....	20
Tableau I.2 Tableau comparatif de Data Profiling.....	21
Tableau I.3Tableau comparatif de Génération de données.....	22

Liste des abréviations

SB	Soft Builder
SGBD	Système de gestion de base de données
SGBDR	Système de gestion de base de données
CSS	Cascading Style Sheet
DOM	Document Object Model
TDM	Toad Data Point
HTML	Hyper Text Markup Language
HTTPS	Hyper Text Transfer Protocol
IHM	Interface Homme Machine
SQL	Structured Query Language
UML	Unified Modeling Language
URL	Uniform Resource Locator
XML	Extensible Markup Language
DTD	Document Type Definition

RÉSUMÉ:

La manipulation des données est une nécessité pour tous les utilisateurs quelle que soit développeurs ou utilisateur simple.

Il est plus facile et plus pratique d'utiliser un outil qui manipule les données pour la gestion des données en accédant à plusieurs SGBD. Dans ce cadre, nous avons développé un système d'administration et de manipulation de données qui permet d'intégrer une bonne partie des fonctionnalités de gestion de données.

Mots Clés:

Base de données, profilage de données, génération de données, auto-complétion, manipulation de données.

ABSTRACT:

The manipulation of data has become something indispensable for developers or simple user.

It's easier and convenient to use a service that manipulate data for data management by accessing several DBMS. In this context, We have developed a data administration and a manipulation system , that allows us to integrate a good part of data management functions.

Keywords:

Data base, data profiling, data generation, auto-completion, data manipulation.

ملخص:

التعامل مع البيانات هو ضرورة لجميع المستخدمين بغض النظر المطورين أو مستخدم بسيط.

بالإضافة إلى أنه من الأسهل والأكثر ملائمة استخدام الخدمة التي تعالج البيانات مع إدارة البيانات عن طريق الوصول إلى العديد من الأنظمة التي تدير قواعد البيانات.

كلمات مفتاحية:

قاعدة البيانات ، احصائيات على البيانات ، توليد البيانات ، التكملة الالية ، معالجة البيانات.