

(19) 日本国特許庁 (JP)

(12) 特 許 公 報 (B2)

(11) 特許番号
特許第6929971号
(P6929971)

(45) 発行日 令和3年9月1日 (2021. 9. 1)

(24) 登録日 令和3年8月13日 (2021. 8. 13)

(51) Int. Cl.

F I

GO 6 N 3/08 (2006. 01)

GO 6 N 3/08

GO 6 F 16/2452 (2019. 01)

GO 6 F 16/2452

請求項の数 19 (全 37 頁)

(21) 出願番号	特願2019-563399 (P2019-563399)	(73) 特許権者	506332063
(86) (22) 出願日	平成30年5月17日 (2018. 5. 17)		セールスフォース ドット コム インコ
(65) 公表番号	特表2020-520516 (P2020-520516A)		ーポレイティッド
(43) 公表日	令和2年7月9日 (2020. 7. 9)		アメリカ合衆国 カリフォルニア州 94
(86) 国際出願番号	PCT/US2018/033099		105, サン フランシスコ, ミッショ
(87) 国際公開番号	W02018/213530		ン ストリート 415, サード フロ
(87) 国際公開日	平成30年11月22日 (2018. 11. 22)		アー
審査請求日	令和1年11月15日 (2019. 11. 15)	(74) 代理人	100107766
(31) 優先権主張番号	15/885, 613		弁理士 伊東 忠重
(32) 優先日	平成30年1月31日 (2018. 1. 31)	(74) 代理人	100070150
(33) 優先権主張国・地域又は機関			弁理士 伊東 忠彦
	米国 (US)	(74) 代理人	100091214
(31) 優先権主張番号	62/508, 367		弁理士 大貫 進介
(32) 優先日	平成29年5月18日 (2017. 5. 18)		
(33) 優先権主張国・地域又は機関			
	米国 (US)		

最終頁に続く

(54) 【発明の名称】 自然言語クエリのデータベースクエリへのニューラルネットワークに基づく翻訳

(57) 【特許請求の範囲】

【請求項 1】

1つ以上のコンピュータを含むコンピュータシステムにより実行される方法であって、前記コンピュータシステムは、

データベーススキーマを使用して記憶されたデータに基づく入力自然言語クエリを受信するステップと、

前記入力自然言語クエリのターム、
前記データベーススキーマの列のセット、及び
データベースクエリ言語の語彙

を含む複数のタームからトークンのシーケンスを生成するステップと、

1つ以上の入力表現を生成するステップであり、各入力表現は、前記トークンのシーケンスをエンコードすることにより取得される、ステップと、

複数の機械学習に基づくモデルにアクセスするステップであり、各々の機械学習に基づくモデルは、前記入力自然言語クエリに対応するデータベースクエリの部分を予測するように構成される、ステップと、

前記複数の機械学習に基づくモデルの各々について、前記機械学習に基づくモデルを入力表現を使用して実行して前記データベースクエリの部分を生成するステップと、

前記データベースクエリの前記生成された部分を組み合わせて前記データベースクエリを取得するステップと、

前記データベースクエリを実行して結果セットを取得するステップと、

10

20

を実行する、方法。

【請求項 2】

前記入力自然言語クエリは、クライアントデバイスから受信され、当該方法は、前記結果セットを前記クライアントデバイスに送信するステップをさらに含む、請求項 1 に記載の方法。

【請求項 3】

前記複数の機械学習に基づくモデルは、前記データベースクエリ内の集約演算子を決定する集約分類器モデルを含み、前記集約分類器モデルは、多層パーセプトロンを含む、請求項 1 に記載の方法。

【請求項 4】

前記複数の機械学習に基づくモデルは、前記データベースクエリの結果列を決定する結果列予測器モデルを含み、前記結果列予測器モデルは、多層パーセプトロンを含む、請求項 1 に記載の方法。

【請求項 5】

前記複数の機械学習に基づくモデルは、前記データベースクエリの条件句を決定する条件句予測器モデルを含み、前記条件句予測器モデルは、強化学習に基づく、請求項 1 に記載の方法。

【請求項 6】

グラウンドトルースデータベースクエリに基づく結果セットを受信するステップであり、前記グラウンドトルースデータベースクエリは、前記入力自然言語クエリに対応し、かつ実行の結果としてグラウンドトルースを提供するクエリを表す、ステップと、

前記生成されたクエリから取得された結果と前記グラウンドトルースデータベースクエリから取得された結果との比較に基づいて報酬値を決定するステップと、

前記報酬値に基づいて前記条件句予測器モデルの重みを調整するステップと、

をさらに含む請求項 5 に記載の方法。

【請求項 7】

前記シーケンスの各トークンに対応する列エンコーディングを決定するステップと、

前記シーケンスの各トークンについてのスコアのベクトルを決定するステップと、

前記ベクトルをソフトマックス関数を使用して正規化するステップと、

前記入力表現を、対応する正規化されたスコアにより重み付けされた前記列エンコーディングの和として決定するステップと、

をさらに含む請求項 5 に記載の方法。

【請求項 8】

前記複数の機械学習に基づくモデルを勾配降下を使用して訓練して、前記複数のモデルの各々の結果に基づいて損失を表す目的関数を最小化するステップ

をさらに含む請求項 1 に記載の方法。

【請求項 9】

前記 1 つ以上の入力表現を生成するステップは、入力表現 κ^{agg} を、

トークンのシーケンス内の各 t 番目のトークンについて、スコア $\alpha^{inp}_t = W^{inp} h^{enc}_t$ を計算し、 h^{enc}_t は、前記シーケンス内の t 番目のワードに対応するエン

前記スコアのベクトル $\alpha^{inp} = [\alpha^{inp}_1, \alpha^{inp}_2, \dots]$ を正規化して、前記トークンのシーケンス内の前記トークンにわたる分布を生成し、

前記入力表現 κ^{agg} を、

【数 1】

$$\kappa^{agg} = \sum_t \beta_t^{inp} h_t^{enc}$$

として取得し、 $\beta^{agg} = \text{softmax}(\alpha^{agg})$ 及び $\alpha^{agg} = W^{agg} \tanh$

$(V^{agg} \cdot K^{agg} + b^{agg}) + c^{agg}$ である

ことにより計算するステップを含む、請求項 1 に記載の方法。

【請求項 10】

1 つ以上の列名を有する前記データベースクエリの select 句を定式化するポインタネットワークで長短期記憶 (LSTM) を使用するステップであり、

列表現のリスト及び質問表現を所与として、質問に最もマッチする列を選択することを含み、前記列表現のリストは、各列名を LSTM でエンコードすることにより取得され、特定の列 j の表現 e_j^c は、

【数 2】

$$h_{j,t}^c = \text{LSTM}(\text{emb}(x_{j,t}^c), h_{j,t-1}^c) \quad e_j^c = h_{j,T_j}^c$$

10

により与えられる、ステップ、

をさらに含む請求項 1 に記載の方法。

【請求項 11】

強化学習を使用して、データベースに対して前記生成されたデータベースクエリを実行することにより前記データベースクエリの where 条件を定式化して、

【数 3】

$$R(q(y), q_g) = \begin{cases} -2, & q(y) \text{ が有効なデータベースクエリでない場合} \\ -1, & q(y) \text{ が有効なデータベースクエリであり、不正確な結果に対し実行する場合} \\ +1, & q(y) \text{ が有効なデータベースクエリであり、正確な結果に対し実行する場合} \end{cases}$$

20

として定義される報酬 $R(q(y), q_g)$ を取得するステップであり、 $q(y)$ は、前記機械学習に基づくモデルにより生成されたクエリを表し、 q_g は、グラウンドトゥースクエリを表し、前記グラウンドトゥースクエリは、前記入力自然言語クエリに対応し、かつ実行の結果としてグラウンドトゥースを提供するクエリを表す、ステップ

をさらに含む請求項 1 に記載の方法。

【請求項 12】

自然言語要求を処理して SQL クエリの部分を生成するステップであり、

1 つ以上の列名を有する前記 SQL クエリの select 句を定式化するポインタネットワークで長短期記憶 (LSTM) を使用することと、

強化学習により訓練されたモデルを使用して前記 SQL クエリの where 条件を定式化することと、を含む、ステップ、

を含む方法。

【請求項 13】

多層パーセプトロンを使用して、指定された条件下で選択された列に適用可能な前記 SQL クエリの集約演算を決定するステップ、

をさらに含む請求項 12 に記載の方法。

40

【請求項 14】

1 つ以上の列名を有する前記 SQL クエリの select 句を定式化するポインタネットワークで長短期記憶 (LSTM) を使用することは、

列表現のリスト及び質問表現を所与として、質問に最もマッチする列を選択することを含む、請求項 12 に記載の方法。

【請求項 15】

前記列表現のリストは、各列名を LSTM でエンコードすることにより取得され、特定の列 j の表現 e_j^c は、

【数 4】

$$h_{j,t}^c = \text{LSTM}(\text{emb}(x_{j,t}^c), h_{j,t-1}^c) \quad e_j^c = h_{j,T_j}^c$$

により与えられる、請求項 1 4 に記載の方法。

【請求項 1 6】

前記 1 つ以上の入力表現を生成することは、入力表現 κ^{agg} を、

トークンのシーケンス内の各 t 番目のトークンについて、スコア $\alpha^{\text{inp}}_t = W^{\text{inp}} h^{\text{enc}}_t$ を計算し、 h^{enc}_t は、前記シーケンス内の t 番目のワードに対応するエンコードの状態であり、

前記スコアのベクトル $\alpha^{\text{inp}} = [\alpha^{\text{inp}}_1, \alpha^{\text{inp}}_2, \dots]$ を正規化して、前記トークンのシーケンス内の前記トークンにわたる分布を生成し、

前記入力表現 κ^{agg} を、

【数 5】

$$\kappa^{\text{agg}} = \sum_t \beta_t^{\text{inp}} h_t^{\text{enc}}$$

として取得し、 $\beta^{\text{agg}} = \text{softmax}(\alpha^{\text{agg}})$ 及び $\alpha^{\text{agg}} = W^{\text{agg}} \tanh(V^{\text{agg}} \kappa^{\text{agg}} + b^{\text{agg}}) + c^{\text{agg}}$ である

ことにより計算することを含む、請求項 1 4 に記載の方法。

【請求項 1 7】

強化学習を使用して、データベースに対して前記生成された SQL クエリを実行することにより前記 SQL クエリの *where* 条件を定式化して、

【数 6】

$$R(q(y), q_g) = \begin{cases} -2, & q(y) \text{ が有効な SQL クエリでない場合} \\ -1, & q(y) \text{ が有効な SQL クエリであり、不正確な結果に対し実行する場合} \\ +1, & q(y) \text{ が有効な SQL クエリであり、正確な結果に対し実行する場合} \end{cases}$$

として定義される報酬 $R(q(y), q_g)$ を取得するステップであり、 $q(y)$ は、機械学習に基づくモデルにより生成されたクエリを表し、 q_g は、グラウンドトゥースクエリを表し、前記グラウンドトゥースクエリは、入力自然言語クエリに対応し、かつ実行の結果としてグラウンドトゥースを提供するクエリを表す、ステップ

をさらに含む請求項 1 4 に記載の方法。

【請求項 1 8】

1 つ以上のコンピュータにステップを実行させるコンピュータプログラムであって、前記ステップは、

データベーススキーマを使用して記憶されたデータに基づく入力自然言語クエリを受信するステップと、

前記入力自然言語クエリのターム、

前記データベーススキーマの列のセット、及び

データベースクエリ言語の語彙

を含む複数のタームからトークンのシーケンスを生成するステップと、

1 つ以上の入力表現を生成するステップであり、各入力表現は、前記トークンのシーケンスをエンコードすることにより取得される、ステップと、

複数の機械学習に基づくモデルにアクセスするステップであり、各々の機械学習に基づくモデルは、前記入力自然言語クエリに対応するデータベースクエリの部分を予測するように構成される、ステップと、

前記複数の機械学習に基づくモデルの各々について、前記機械学習に基づくモデルを入力表現を使用して実行して前記データベースクエリの部分を生成するステップと、

10

20

30

40

50

前記データベースクエリの前記生成された部分を組み合わせて前記データベースクエリを取得するステップと、

前記データベースクエリを実行して結果セットを取得するステップと、
を含む、コンピュータプログラム。

【請求項 19】

リカレントニューラルネットワーク (R N N) アーキテクチャを生成するコンピュータシステムであって、

1 つ以上のコンピュータプロセッサと、

1 つ以上のコンピュータにステップを実行させるコンピュータプログラムと、

を含み、前記ステップは、

データベーススキーマを使用して記憶されたデータに基づく入力自然言語クエリを受信するステップと、

前記入力自然言語クエリのターム、

前記データベーススキーマの列のセット、及び

データベースクエリ言語の語彙

を含む複数のタームからトークンのシーケンスを生成するステップと、

1 つ以上の入力表現を生成するステップであり、各入力表現は、前記トークンのシーケンスをエンコードすることにより取得される、ステップと、

複数の機械学習に基づくモデルにアクセスするステップであり、各々の機械学習に基づくモデルは、前記入力自然言語クエリに対応するデータベースクエリの部分を予測するように構成される、ステップと、

前記複数の機械学習に基づくモデルの各々について、前記機械学習に基づくモデルを入力表現を使用して実行して前記データベースクエリの部分を生成するステップと、

前記データベースクエリの前記生成された部分を組み合わせて前記データベースクエリを取得するステップと、

前記データベースクエリを実行して結果セットを取得するステップと、

を含む、コンピュータシステム。

【発明の詳細な説明】

【技術分野】

【0001】

本開示は、一般に、データベースクエリの自動生成に関し、より具体的には、自然言語クエリをデータベースクエリに翻訳するためのニューラルネットワークに基づくモデルに関する。

【背景技術】

【0002】

世界中で利用可能な相当量のデータは、関係データベースに記憶されている。関係データベースは、医療記録、金融市場、顧客関係管理などのアプリケーションの基礎を提供する。しかしながら、関係データベース内の情報へのアクセスは、構造化問い合わせ言語 (S Q L) などのデータベースクエリ言語の理解を要する。S Q L などのデータベースクエリ言語は、ユーザが関係データベースからのデータの要求を指定することを可能にする点で強力であるが、これらは学習することが困難である。データベースクエリ言語を使用して効率的にデータベースクエリを書けるようにするには、データベースにおける専門知識と強力な技術知識を要する。

【0003】

一部のシステムは、システムに記憶されたデータにアクセスするための自然言語をサポートしている。自然言語クエリは、人々が自然言語を使用するための訓練を必要としないため、表現の容易さを提供する。しかしながら、これらのシステムは、S Q L などのデータベースクエリ言語の表現力を提供しない。例えば、自然言語クエリは、複数の方法で解釈される可能性があり、関係データベースに記憶されたデータにアクセスするための自然言語クエリの対応する実行は、非効率的な可能性があり、要求された正確な情報を取り出

10

20

30

40

50

さない可能性がある。したがって、自然言語クエリ又はデータベースクエリのいずれかを使用して関係データベースに記憶されたデータにアクセスする従来の手法は、これらが表現の容易さを提供するか又は表現の力を提供するかのいずれかであり、あるいは双方を提供しないため、欠点を有する。

【図面の簡単な説明】

【0004】

開示される実施形態は、詳細な説明、別記の特許請求の範囲、及び添付の図（又は図面）からより容易に明らかとなる他の利点及び特徴を有する。以下は、図の簡単な紹介である。

【図1】一実施形態による、自然言語クエリをデータベースクエリに翻訳するための全体的なシステム環境を示す高レベルブロック図である。

10

【図2】一実施形態による、自然言語クエリをデータベースクエリに翻訳するためのコンピューティングシステムのシステムアーキテクチャを示す。

【図3】一実施形態による、自然言語対データベースクエリ翻訳器により実行される処理の詳細を示す。

【図4】一実施形態による、自然言語クエリをデータベースクエリに翻訳するための全体的な処理を示す。

【図5】一実施形態による、自然言語クエリに基づいて出力データベースクエリの集約演算子を決定する集約分類器の処理を示す。

【図6】一実施形態による、自然言語クエリに基づいて出力データベースクエリのSELECT句の列を決定する結果列予測器の処理を示す。

20

【図7】一実施形態による、出力データベースクエリの条件句を決定する条件句予測器を訓練する処理を示す。

【図8】図1のクライアントデバイス及び/又はコンピューティングシステムを実現するための一例示的なコンピュータを示す高レベルブロック図である。

【0005】

図（図面）及び以下の説明は、特定の実施形態を単なる例示として説明する。当業者は、以下の説明から、本明細書に示された構造及び方法の代替的な実施形態が、本明細書で説明される原理から逸脱することなく採用され得ることを容易に理解するであろう。次に、いくつかの実施形態を詳細に参照し、その例を添付の図に示す。

30

【発明を実施するための形態】

【0006】

コンピューティングシステムが、自然言語クエリ（natural language queries）に対応するデータベースクエリ、例えば構造化問い合わせ言語（structured query language、SQL）を使用して指定されたクエリに翻訳するために、ディープニューラルネットワークを使用する。実施形態は、SQLクエリの構造を使用して、生成されたクエリの出力空間を大幅に削減する。コンピューティングシステムは、ディープニューラルネットワークを使用して、自然言語クエリをデータベースクエリに翻訳する。

【0007】

一実施形態において、コンピューティングシステムは、複数の機械学習に基づくモデル、例えばニューラルネットワークに基づくモデルを使用して、出力データベースクエリの異なる部分を生成する。例えば、コンピューティングシステムは、データベースクエリ内の集約演算子を決定する集約分類器（aggregation classifier）モデル、データベースクエリの結果列を決定する結果列予測器（result column predictor）モデル、及びデータベースクエリの条件句を決定する条件句予測器（condition clause predictor）モデルを使用することができる。一実施形態において、集約分類器モデル及び結果列予測器モデルは、多層パーセプトロンを含む。条件句予測器モデルは、ポリシーに基づく強化学習（reinforcement learning、RL）を使用して、データベースクエリの条件句を生成する。これは、条件句が本質的に順序づけられておらず、条件句の複数の表現がデータベースクエリに対して同じ出力結果を提供する可能性があるためである。ゆえに、条件句は

40

50

、交差エントロピー損失 (cross entropy loss) を使用した最適化に適さない。ディープニューラルネットワークは、交差エントロピー損失と R L 報酬とを組み合わせた混合された目的 (mixed objective) を使用して訓練される。

【 0 0 0 8 】

一例として、データベースが、列 Pick_number、CFL_Team、Player、Position、及び College を有するテーブル CFL Draft を記憶することができる。このテーブルは、以下の例示的な行を記憶することができる。

【表 1】

Pick_number	CFL_Team	Player	Position	College
27	Hamilton Tiger-Cats	Connor Healy	DB	Wilfrid Laurier
28	Calgary Stampeders	Anthony Forgone	OL	York
29	Ottawa Renegades	L. P. Ladouceur	DT	California
30	Toronto Argonauts	Frank Hoffman	DL	York
...	...	...	...	...

10

20

【 0 0 0 9 】

システムは、自然言語クエリ、例えば「いくつの CFL チームがヨーク大学からか? (How many CFL teams are from York College?)」を受信する。システムは、テーブル CFL Draft を含むデータベーススキーマに関連して受信した自然言語クエリを処理して、SQL 言語を使用したデータベースクエリ「SELECT COUNT (CFL_Team) FROM CFL Draft WHERE College = "York"」を生成する。システムは、データベーススキーマを使用してデータベースクエリを実行する。テーブル CFL Draft の 2 行が、これらが大学「York」を有するため、データベースクエリの WHERE 句にマッチする。結果として、システムは結果 2 を返す。

30

【 0 0 1 0 】

[全体的なシステム環境]

図 1 は、一実施形態による、自然言語クエリをデータベースクエリに翻訳するための全体的なシステム環境を示す高レベルブロック図である。システム環境 100 は、ネットワーク 150 によりコンピューティングシステム 130 に接続された 1 つ以上のクライアントデバイス 110 を含む。コンピューティングシステム 130 は、オンラインシステムであってよいが、例えば、自然言語クエリのセットの各々をデータベースクエリに翻訳するバッチ処理を実行することによりオフラインで動作してもよい。

40

【 0 0 1 1 】

2 つのクライアントデバイス 110 a、110 b のみがここで示されているが、これらのエンティティの各々の、複数のインスタンスが存在してもよい。例えば、いくつかのコンピューティングシステム 130 と、各コンピューティングシステム 130 と通信する数十又は数百のクライアントデバイス 110 とが存在してもよい。図面は、同様の要素を識別するために、同様の参照番号を使用する。参照番号の後の文字、「110 a」などは、テキストがその特定の参照番号を有する要素を具体的に参照することを示す。続きの文字のないテキストの参照番号、「110」などは、その参照番号を有する図中の要素のい

50

れか又は全てを参照する。

【0012】

クライアントデバイス110は、ANDROID（登録商標）若しくはAPPLE（登録商標）IOS（登録商標）などのオペレーティングシステムを有するスマートフォン、タブレットコンピュータ、ラップトップコンピュータ、デスクトップコンピュータ、自動車若しくは他の車両の電子ステレオ、又はデジタルコンテンツが聴けるか又は他の方法で体験できる任意の他タイプのネットワーク対応デバイスなどのコンピューティングデバイスである。典型的なクライアントデバイス110は、（例えば、Wifi及び/又は4G又は他の無線電気通信標準を介して）ネットワーク150に接続するために必要なハードウェア及びソフトウェアを含む。

10

【0013】

クライアントデバイス110は、クライアントデバイス110のユーザがコンピューティングシステム130と対話することを可能にするクライアントアプリケーション120を含む。例えば、クライアントアプリケーション120は、ユーザがコンピューティングシステム130に送信される自然言語クエリを入力することを可能にするユーザインターフェースであってもよい。クライアントアプリケーション120は、コンピューティングシステム130から結果を受信し、これらをユーザインターフェースを介してユーザに提示する。一実施形態において、クライアントアプリケーション120は、クライアントデバイス110のユーザがコンピューティングシステム130上で実行しているウェブサーバと対話することを可能にするブラウザである。

20

【0014】

コンピューティングシステム130は、協調された機能又はタスクのグループを実行するソフトウェアを含む。ソフトウェアは、コンピューティングシステム130のユーザが関心のある特定のタスク又はアクティビティを実行することを可能にでき、あるいは他のソフトウェアに特定の機能性及びサービスを提供するシステムソフトウェア（例えば、オペレーティングシステム）を含んでもよい。コンピューティングシステム130は、クライアントデバイス110から要求を受信し、受信した要求に関連づけられたコンピュータプログラムを実行する。一例として、コンピューティングシステム130は、クライアントデバイス110からの、自然言語クエリをデータベースクエリに翻訳する要求に回答して、コンピュータプログラムを実行してもよい。コンピューティングシステム130上で実行するソフトウェアは、複数の当事者又はチームがソフトウェアの異なるコンポーネントを管理する責任を負う協調的な方法で書かれたコンピュータプログラム、ライブラリ、及び関連データの複雑なコレクションを含むことができる。

30

【0015】

一実施形態において、コンピューティングシステム130は、クライアントデバイス110から自然言語クエリ135を受信する。自然言語クエリ135は、コンピューティングシステム130上で実行するクライアントアプリケーション120を介してユーザにより提供されてもよい。コンピューティングシステム130は、データベースに記憶されるデータの構造を定義するデータベーススキーマ145を記憶する。例えば、データベーススキーマ145は、データベースに記憶された様々なテーブル、各テーブルの列、外部キー関係などのテーブル間の関係、テーブルに関連づけられた任意の制約などを識別することができる。

40

【0016】

自然言語対データベースクエリ翻訳器（natural language to database query translator）140は、入力として自然言語クエリ135及びデータベーススキーマ145を受信し、入力された自然言語クエリ135に相当するデータベースクエリ155を生成する。生成されたデータベースクエリ155は、データベーススキーマ145に準拠する。生成されたデータベースクエリ155は、データベースクエリプロセッサ150により受信され、データベースクエリプロセッサ150は、データベース160に記憶されたデータを使用してデータベースクエリ155を処理する。データベースクエリプロセッサ1

50

50は、データベースクエリ155を処理することによりクエリ結果165を生成する。コンピューティングシステム130は、生成されたクエリ結果165を、自然言語クエリ135を送信したクライアントデバイス110上で実行しているクライアントアプリケーション120に提供する。

【0017】

一実施形態において、自然言語対データベースクエリ翻訳器140は、シーケンス対シーケンス(sequence to sequence)翻訳を実行する。従来のニューラルネットワークに基づくシーケンス対シーケンス翻訳器は、かなり大きい空間で検索する。対照的に、実施形態は、データベースクエリ言語に固有の構造を活用して検索空間を削減する。詳細には、システムは、テーブルスキーマ、入力質問、及びSQLキーワードの結合に基づいて、生成されたシーケンスの出力空間を制限する。一実施形態において、自然言語対データベースクエリ翻訳器140は、拡大された(augmented)入力を有するポインタネットワークであるディープニューラルネットワークを使用する。

10

【0018】

ネットワーク150は、クライアントデバイス110とレコード管理システム130との間の通信インフラストラクチャを提供する。ネットワーク150は、典型的にはインターネットであるが、これらに限られないがローカルエリアネットワーク(LAN)、メトロポリタンエリアネットワーク(MAN)、ワイドエリアネットワーク(WAN)、モバイル有線若しくは無線ネットワーク、プライベートネットワーク、又は仮想プライベートネットワークを含む任意のネットワークであってもよい。ネットワーク150の一部が、IEEE802.11標準に基づくWi-Fi、BLUETOOTH(登録商標)ショートレンジ標準、及びワイヤレスユニバーサルシリアルバス(USB)標準を含む通信技術を使用するリンクにより提供されてもよい。

20

【0019】

[システムアーキテクチャ]

図2は、一実施形態による、自然言語クエリをデータベースクエリに翻訳するためのコンピューティングシステムのシステムアーキテクチャを示す。コンピューティングシステム130は、入力エンコーディングモジュール210、訓練モジュール240、自然言語対データベースクエリ翻訳器140、クエリ合成モジュール220、クエリ実行エンジン230、訓練データストア215、及びデータベース160を含む。ネットワークインターフェース、セキュリティ機能、ロードバランサ、フェイルオーバーサーバ、管理及びネットワーク操作コンソールなどの従来のコンポーネントは、システムアーキテクチャの詳細を分かりにくくしないよう示されていない。

30

【0020】

入力前処理モジュール210は、自然言語対データベースクエリ翻訳器140への入力として提供するために入力データを前処理する。一実施形態において、入力前処理モジュール210は、データベーススキーマからの列名と、入力自然言語クエリと、データベースクエリ言語、例えばSQLの語彙(vocabulary)とを連結することにより、トークンのシーケンスを生成する(420)。入力前処理モジュール210は、出力データベースクエリの様々な部分を生成する様々なモデルに提供するために1つ以上の入力表現を生成する。

40

【0021】

自然言語対データベースクエリ翻訳器140は、自然言語クエリに対応するデータベースクエリを生成するために入力自然言語クエリを処理する。一実施形態において、自然言語対データベースクエリ翻訳器140は、図3に関連して本明細書でさらに説明される他のコンポーネント、例えば、集約分類器260、結果列予測器270、及び条件句予測器280を含む。

【0022】

自然言語対データベースクエリ翻訳器140は、異なるニューラルネットワークを使用してデータベースクエリの異なるコンポーネントを生成する。一実施形態において、自然

50

言語対データベースクエリ翻訳器 140 は、異なるニューラルネットワークを使用して、選択列と集約演算子と where 句とを含むデータベースクエリのコンポーネントを生成する。

【0023】

訓練モジュール 240 は、訓練データストア 215 に記憶された過去データを使用して、自然言語対データベースクエリ翻訳器 140 内のニューラルネットワークを訓練する。一実施形態において、訓練モジュール 240 は、集約分類器 260 及び結果列予測器 270 を交差エントロピー損失を使用して訓練するが、条件句予測器 280 をポリシー勾配 (policy gradient) 強化学習を使用して訓練して、クエリ条件の順序づけられていない性質に対処する。SQL クエリの構造を利用することで、自然言語対データベースクエリ翻訳器 140 がデータベースクエリの出力空間を削減することを可能にする。これは、クエリ構造を活用しない他の手法と比較して、有意により高い性能をもたらす。

10

【0024】

クエリ合成モジュール 220 は、自然言語対データベースクエリ翻訳器 140 により生成されたデータベースクエリの様々なコンポーネントを受信し、これらを組み合わせてデータベースクエリを取得する。クエリ実行モジュール 230 は、データベース 160 に記憶されたデータを使用して、クエリ合成モジュール 220 により提供されたデータベースクエリを実行する。コンピューティングシステム 130 は、クエリの実行の結果を、結果の要求元、例えばクライアントデバイス 110 上で実行しているクライアントアプリケーション 120 に返す。

20

【0025】

図 3 は、一実施形態による、自然言語対データベースクエリ翻訳器 140 により実行される処理の詳細を示す。図 3 に示すように、自然言語対データベースクエリ翻訳器 140 への入力、自然言語クエリ 320 及びデータベーススキーマ 320 を含む。CFLDraft テーブルに基づく上述の例では、自然言語クエリ 320 は、「いくつかの CFL チームがヨーク大学からか？」であり、データベーススキーマ 320 は、列 Pick_number、CFL_Team、Player、Position、及び College を含む様々な列を含む。例示的な出力データベースクエリは、「SELECT COUNT (CFL_Team) FROM CFLDraft WHERE College = "York"」である。

30

【0026】

入力前処理モジュール 210 は、1 つ以上の入力表現を生成し、集約分類器 260、結果列予測器 270、及び条件句予測器 280 を含む自然言語対データベースクエリ翻訳器 140 の各コンポーネントに入力表現を提供する。集約分類器 260、結果列予測器 270、及び条件句予測器 280 の各々は、出力データベースクエリの一部を生成する。

【0027】

結果列予測器 270 は、結果列、例えば、SQL を使用して表現された出力データベースクエリの SELECT 句 310 で指定される列を生成する。結果列の一例が、例示的な出力データベースクエリ内の列 CFL_Team である。一実施形態において、結果列予測器 270 は、列のシーケンスのエンコーディングを入力として受信し、かつ SELECT 列に対応する該列のシーケンス内の列を指し示すポインタネットワークである。

40

【0028】

条件句予測器 280 は、出力データベースクエリの出力行をフィルタリングするために使用される条件を指定する出力データベースクエリの WHERE 句 320 を生成する。上記の例では、WHERE 句「College = "York"」が出力データベースクエリ内の条件句である。

【0029】

集約分類器 260 は、もしあれば出力データベースクエリ内の集約演算子 330、例えば、例示的な出力データベースクエリ内の COUNT 演算子を生成する。集約演算子は、SQL により選択された行の要約を生成する。集約分類器 260 により生成され得る集約

50

演算子の例は、最大 (MAX)、最小 (MIN)、平均 (AVG)、和 (SUM) などを含む。集約分類器 260 は、出力クエリに集約演算子が存在しない場合、NULL 集約演算子を生成することができる。

【0030】

SELECT 句 310、WHERE 句 320、及び集約演算子 330 を含む出力データベースクエリの様々なコンポーネントは、クエリ合成モジュール 270 に入力として提供される。クエリ合成モジュール 270 は、出力データベースクエリの個々のコンポーネントを組み合わせて、完全な出力データベースクエリ 340 を生成する。

【0031】

[全体的な処理]

図 4 ~ 図 7 は、自然言語クエリをデータベースクエリに翻訳するための様々な処理を示す。当業者は、他の実施形態が図 4 ~ 図 7 のステップをフローチャートに示される順序と異なる順序で実行できることを認識するであろう。さらに、他の実施形態は、本明細書で説明されるステップと異なる及び / 又は追加のステップを含むことができる。特定のモジュールにより実行されるものとして示されるステップは、他のモジュールにより実行されてもよい。

【0032】

図 4 は、一実施形態による、自然言語クエリをデータベースクエリに翻訳するための全体的な処理を示す。自然言語対データベースクエリ翻訳器 140 が、入力自然言語クエリを受信する (410)。入力前処理モジュール 210 が、データベーススキーマからの列名と、入力自然言語クエリと、データベースクエリ言語の語彙、例えば、SELECT、FROM、WHERE などの SQL 言語の様々なキーワードとを連結することにより、トークンのシーケンスを生成する (420)。例えば、式 (1) は、列名 x_i^c と、SQL 語彙を表すターム (terms) x^s と、入力自然言語クエリを表すターム x^q とを含むトークンのシーケンスを示す。

【数 1】

$$x = [< col >; x_1^c; x_2^c; \dots; x_N^c; < sql >; x^s; < question >; x^q] \quad (1)$$

【0033】

式 (1) において、シーケンス a 及び b の間の連結は [a ; b] で表される。さらに、組み合わせられたシーケンス x は、境界を画定するために、近隣のシーケンスの間にセンチネルトークンを含む。例えば、トークン < col > は列名を識別し、トークン < sql > は SQL 語彙を表すタームを識別し、トークン < question > は入力自然言語クエリのタームを識別する。

【0034】

入力前処理モジュール 210 は、トークンのシーケンスの入力表現を生成する (430)。一実施形態において、入力前処理モジュール 210 は、複数のモデルの各々について 1 つで、複数の入力表現を生成する。

【0035】

自然言語対データベースクエリ翻訳器 140 は、複数のニューラル機械学習モデルにアクセスし、各モデルは、出力データベースクエリの部分を生成するように構成される。一実施形態において、自然言語対データベースクエリ翻訳器 140 は、複数の訓練されたニューラルネットワークに基づくモデルをストレージデバイスからメモリにロードする。自然言語対データベースクエリ翻訳器 140 は、複数の機械学習に基づくモデルの各々に入力表現を提供する (450)。複数の機械学習に基づくモデルの各々は、データベースクエリの部分を生成する。

【0036】

いくつかの実施形態において、入力前処理モジュール 210 は、複数の入力表現を生成し、自然言語対データベースクエリ翻訳器 140 は、各々の機械学習に基づくモデルに異なる入力表現を提供することができる。各々の機械学習に基づくモデルは、データベース

10

20

30

40

50

クエリの部分を生成し、それをクエリ合成モジュール 270 に提供する。クエリ合成モジュール 270 は、データベースクエリの複数の部分を組み合わせてフルのデータベースクエリを生成する (460)。クエリ実行エンジン 230 が、データベースクエリを実行して結果セットを生成する (470)。

【0037】

〔集約分類器〕

図 5 は、一実施形態による、自然言語クエリに基づいて出力データベースクエリの集約演算子を決定する集約分類器の処理を示す。集約分類器 260 は、入力自然言語クエリで指定された質問のタイプに基づいて、出力データベースクエリの集約演算子を決定する。例えば、集約分類器 260 は、文字列「いくつの (how many)」を含む入力質問を集約演算子 COUNT にマッピングすることができ、集約分類器 260 は、「何が最も高い」を含む入力質問を集約演算子最大にマッピングすることができ、集約分類器 260 は、「何が最も小さい」を含む入力質問を集約演算子最小にマッピングすることができる、などである。

10

【0038】

集約分類器 260 は、トークンの入力シーケンスの入力表現を決定する (510)。集約分類器 260 は、入力シーケンス内の各 t 番目のトークンについて、スカラー注目スコア (scalar attention score) $\alpha_t^{inp} = W^{inp} * h_t^{enc}$ を計算する。したがって、集約分類器 260 は、スコアのベクトル $\alpha^{inp} = [\alpha_1^{inp}, \alpha_2^{inp}, \dots]$ を生成する。集約分類器 260 は、 α^{inp} ベクトルにソフトマックス関数を適用して $\beta^{inp} = \text{softmax}(\alpha^{inp})$ を決定することにより、スコアのベクトル α^{inp} を正規化し、入力エンコーディングにわたる分布を生成する。集約分類器 260 は、入力エンコーディングにわたる分布を生成する。集約分類器 260 は、入力表現 κ^{agg} を、以下の式で示されるように、正規化されたスコア β^{inp} により重み付けされた入力エンコーディング h^{enc} にわたる和として決定する (510)。

20

【数 2】

$$\kappa^{agg} = \sum_t \beta_t^{inp} h_t^{enc} \quad (2)$$

【0039】

集約分類器 260 は、生成された入力表現 κ^{agg} に適用される多層パーセプトロンを含み、例えば、COUNT、MIN、MAX、集約なしを示す NULL 演算子など、様々な集約演算に対応するスコア α^{agg} を生成する。集約分類器 260 は、生成されたスコアに基づいて、データベースクエリに対する集約演算を識別する (530)。

30

【0040】

一実施形態において、集約分類器 260 は、以下の式を使用して α^{agg} を決定する。

【数 3】

$$\alpha^{agg} = W^{agg} \tanh(V^{agg} \kappa^{agg} + b^{agg}) + c^{agg} \quad (3)$$

40

【0041】

項 W^{agg} 、 V^{agg} 、 b^{agg} 及び c^{agg} は、多層パーセプトロンに対応する重みを表す。集約分類器 260 は、ソフトマックス関数を適用して、可能な集約演算のセットにわたる分布 $\beta^{agg} = \text{softmax}(\alpha^{agg})$ を取得する。集約分類器は、交差エントロピー損失 L^{agg} に基づいて訓練される。

【0042】

〔結果列予測器〕

SELECT 句は、選択列又は結果列とも呼ばれる。結果列予測器 270 は、データベーススキーマ内のテーブル列及び自然言語クエリに基づいて選択列を決定する。例えば、自然言語クエリ「いくつの C F L チームが・・・」を所与として、結果列予測器 270 は

50

、選択列が C F L D r a f t テーブルからの C F L _ T e a m s 列を含むことを決定する。したがって、結果列予測器 270 は、S E L E C T 列予測の問題をマッチング問題として解決する。一実施形態において、結果列予測器 270 は、ポインタを使用して S E L E C T 列を識別する。列表現のリスト及び自然言語クエリの表現を所与として、結果列予測器 270 は、自然言語クエリに最もマッチする列を選択する。

【0043】

図6は、一実施形態による、自然言語クエリに基づいて出力データベースクエリの S E L E C T 句の列を決定する結果列予測器により実行される処理を示す。結果列予測器 270 は、各列名を L S T M (長短期記憶ネットワーク (long short term memory network)) でエンコードすること (610) により、列に対する入力表現を使用する。入力前

10

【数4】

$$h_{j,t}^c = \text{LSTM}(\text{emb}(x_{j,t}^c), h_{j,t-1}^c) \quad e_j^c = h_{j,Tj}^c \quad (4)$$

【0044】

この式において、 $h_{j,t}^c$ は j 番目の列の t 番目のエンコード状態を表し、 emb は埋め込みを返す関数である。入力前処理モジュール 210 は、最後のエンコード状態を、 e_j^c 、列 j の表現であると見なす。

20

【0045】

入力前処理モジュール 210 は、 κ^{agg} について上述したものと同様のアーキテクチャを使用して、自然言語クエリ κ^{sel} に対する表現を構築する。結果列予測器 270 は、入力表現を条件として列表現にわたり多層パーセプトロンを適用して、以下の式を使用して各列 j についてスコアを計算する (630)。

【数5】

$$\alpha_j^{sel} = W^{sel} \tanh(V^{sel} \kappa^{sel} + V^c e_j^c) \quad (5)$$

【0046】

30

この式において、 W^{sel} 、 V^{sel} 、及び V^c は、多層パーセプトロンの重みである。結果列予測器 270 は、スコアをソフトマックス関数で正規化して、可能な S E L E C T 列にわたる分布 $\beta^{sel} = \text{softmax}(\alpha^{sel})$ を生成する (640)。上記の C F L D r a f t テーブルの例では、分布は、列 P i c k _ n u m b e r、C F L _ T e a m、P l a y e r、P o s i t i o n、及び C o l l e g e にわたる。結果列予測器 270 は、正規化されたスコアに基づいて、出力データベースクエリの結果列を選択する (650)。集約分類器は、交差エントロピー損失 L^{sel} に基づいて訓練される。

【0047】

[条件句予測器]

一実施形態において、条件句予測器は、ポインタデコーダを使用して W H E R E 句を生成する。しかしながら、クエリの W H E R E 条件は入れ替えられる可能性があり、クエリは同じ結果を生成する。例えば、自然言語クエリ「どの男性が18歳より上か (which males are older than 18)」を所与として、出力データベースクエリは、「S E L E C T name FROM insurance WHERE age > 18 AND gender = "male"」又は「S E L E C T name FROM insurance WHERE gender = "male" AND age > 18」のいずれかであり得る。2つのデータベースクエリが、2つのクエリ文字列間の文字列マッチに基づきマッチしない場合でも、双方のデータベースクエリが正しい実行結果を取得する。第1のデータベースクエリが、ニューラルネットワークを訓練する間にグラウンドトゥースとして提供され、交差エントロピー損失が、訓練を監督する (supervise) ために使用され

40

50

る場合、第2のデータベースクエリは、それが文字列マッチに基づき第1のデータベースクエリにマッチしないため、誤ってペナルティを課される。ゆえに、実施形態は、強化学習を適用してポリシーを学習し、データベースクエリの実行結果の期待された正確さを直接最適化する。

【0048】

図7は、一実施形態による、出力データベースクエリの条件句を決定する条件句予測器を訓練する処理を示す。条件句予測器280は、入力として自然言語クエリ710及びデータベーススキーマ720を受信して、データベースクエリ730を生成する。条件句予測器280は、データベース160を使用した実行のためにデータベースクエリを送信して、報酬メトリックを取得する。クエリ実行エンジン230は、生成されたデータベースクエリ730を実行して、予測クエリ結果750を取得する。コンピューティングシステム130は、グラウンドトルースクエリ結果750を訓練データストア215に記憶する。条件句予測器280は、予測クエリ結果750をグラウンドトルースクエリ結果750と比較して、報酬750を決定する。報酬は、条件句予測器280を訓練するためのフィードバックとして、条件句予測器280に入力として提供される。

10

【0049】

WHERE句内の条件句予測器280により生成されるトークンのシーケンスは、 $y = [y^1, y^2, \dots, y^T]$ で表される。 $q(y)$ がモデルにより生成されたクエリを表し、 q_g が自然言語クエリに対応するグラウンドトルースデータベースクエリを表すとする。条件句予測器280は、報酬メトリック $R(q(y), q_g)$ として以下の式を使用する。

20

【数6】

$$R(q(y), q_g) = \begin{cases} -2, & q(y) \text{ が有効なSQLクエリでない場合} \\ -1, & q(y) \text{ が有効なSQLクエリであり、不正確な結果に対し実行する場合} \\ +1, & q(y) \text{ が有効なSQLクエリであり、正確な結果に対し実行する場合} \end{cases} \quad (6)$$

【0050】

したがって、条件句予測器280は、生成されたデータベースクエリの実行の結果がグラウンドトルースとして提供される期待された結果にマッチする場合、正の報酬を割り当てる。条件句予測器280は、生成されたデータベースクエリの実行の結果がグラウンドトルースとして提供される期待された結果にマッチするのに失敗した場合、又は生成されたデータベースクエリが有効なデータベースクエリでない場合、負の報酬を割り当てる。

30

【0051】

条件句予測器280は、損失 $L^{w h e}$ を、可能なWHERE句にわたる負の期待された報酬として決定する。訓練モジュールは、勾配降下 (gradient descent) を使用して条件句予測器280を訓練して、目的関数 $L = L^{a g g} + L^{s e l} + L^{w h e}$ を最小化する。したがって、条件句予測器280は、SELECT列を予測する際の交差エントロピー損失からの、集約演算を予測する際の交差エントロピー損失からの、及び条件句のためのポリシー学習からの勾配の重み付け和として、総勾配を決定する。

【0052】

40

自然言語対データベースクエリ翻訳器140における構造の組み込みは、生成される可能性のある無効なデータベースクエリを削減する。大量の無効なクエリは、列名から結果として生じ、生成されたクエリは、テーブルに存在しない選択列を参照する。これは、列名が多くトークン、例えば4つのトークンを有する「マイル(km)」などを含むとき、特に役立つ。集約のために分類器を導入することも、誤り率を削減する。集約分類器の使用は、COUNT演算子を予測するための適合率及び再現率を向上させる。条件句を生成するための表現学習の使用は、グラウンドトルースと異なって順序づけられ得る、より高品質のWHERE句の生成を結果としてもたらす。ポリシーに基づく表現学習での訓練は、条件の順序がグラウンドトルースクエリと異なる場合でも、正しい結果をもたらす。

【0053】

50

[コンピュータアーキテクチャ]

図8は、図1のクライアントデバイス及び/又はコンピューティングシステムを実現するための一例示的なコンピュータを示す高レベルブロック図である。コンピュータ800は、チップセット804に結合された少なくとも1つのプロセッサ802を含む。チップセット804は、メモリコントローラハブ820及び入力/出力(I/O)コントローラハブ822を含む。メモリ806及びグラフィックスアダプタ812がメモリコントローラハブ820に結合され、ディスプレイ818がグラフィックスアダプタ812に結合される。ストレージデバイス808、入力デバイス814、及びネットワークアダプタ816が、I/Oコントローラハブ822に結合される。コンピュータ800の他の実施形態は、異なるアーキテクチャを有する。

10

【 0054 】

ストレージデバイス808は、ハードドライブ、コンパクトディスク読取専用メモリ(CD-ROM)、DVD、又はソリッドステートメモリデバイスなどの、非一時的コンピュータ読取可能記憶媒体である。メモリ806は、プロセッサ802により使用される命令及びデータを保持する。入力インターフェース814は、タッチスクリーンインターフェース、マウス、トラックボール、若しくは他タイプのポインティングデバイス、キーボード、又はこれらの何らかの組み合わせであり、コンピュータ800にデータを入力するために使用される。いくつかの実施形態において、コンピュータ800は、ユーザからのジェスチャを介して入力インターフェース814から入力(例えば、コマンド)を受け取るように構成されてもよい。グラフィックスアダプタ812は、画像及び他の情報をディスプレイ818に表示する。ネットワークアダプタ816は、コンピュータ800を1つ以上のコンピュータネットワークに結合する。

20

【 0055 】

コンピュータ800は、本明細書で説明される機能性を提供するコンピュータプログラムモジュールを実行するように適合される。本明細書で用いられるとき、用語「モジュール」は、指定された機能性を提供するために使用されるコンピュータプログラム論理を参照する。ゆえに、モジュールは、ハードウェア、ファームウェア、及び/又はソフトウェアで実現できる。一実施形態において、プログラムモジュールは、ストレージデバイス808に記憶され、メモリ806にロードされ、プロセッサ802により実行される。

【 0056 】

図1のエンティティにより使用されるコンピュータ800のタイプは、実施形態及びエンティティにより必要とされる処理能力に依存して変わってもよい。コンピュータ800は、グラフィックスアダプタ812及びディスプレイ818など、上述のコンポーネントのうちいくつかがなくともよい。例えば、コンピューティングシステム130は、サーバファーム内などのネットワークを通して通信する複数のブレードサーバから形成されてもよい。

30

【 0057 】

本出願の対象事項の実現は、特に、以下の例1~21であってもよい。

[例1]

データベーススキーマを使用して記憶されたデータに基づく入力自然言語クエリを受信するステップと、

40

上記入力自然言語クエリのターム、

上記データベーススキーマの列のセット、及び

データベースクエリ言語の語彙

を含む複数のタームからトークンのシーケンスを生成するステップと、

1つ以上の入力表現を生成するステップであり、各入力表現は、上記トークンのシーケンスをエンコードすることにより取得される、ステップと、

複数の機械学習に基づくモデルにアクセスするステップであり、各モデルは、上記入力自然言語クエリに対応するデータベースクエリの部分を予測するように構成される、ステップと、

50

上記複数のモデルの各々について、上記モデルを入力表現に基づいて実行して上記データベースクエリの部分を生成するステップと、

上記データベースクエリの上記生成された部分を組み合わせて上記データベースクエリを取得するステップと、

上記データベースクエリを実行して結果セットを取得するステップと、
を含む方法。

〔例 2〕

上記入力自然言語クエリは、クライアントデバイスから受信され、当該方法は、上記結果セットを上記クライアントデバイスに送信するステップをさらに含む、例 1 に記載の方法。

10

〔例 3〕

上記複数のモデルは、上記データベースクエリ内の集約演算子を決定する集約分類器モデルを含み、上記集約分類器モデルは、多層パーセプトロンを含む、例 1 に記載の方法。

〔例 4〕

上記複数のモデルは、上記データベースクエリの結果列を決定する結果列予測器モデルを含み、上記結果列予測器モデルは、多層パーセプトロンを含む、例 1 に記載の方法。

〔例 5〕

上記複数のモデルは、上記データベースクエリの条件句を決定する条件句予測器モデルを含み、上記条件句予測器モデルは、強化学習に基づく、例 1 に記載の方法。

〔例 6〕

20

グラウンドトゥースデータベースクエリに基づく結果セットを受信するステップと、
上記生成されたクエリから取得された結果と上記グラウンドトゥースクエリから取得された結果との比較に基づいて報酬値を決定するステップと、

上記報酬値に基づいて上記条件句予測器モデルの重みを調整するステップと、
をさらに含む例 5 に記載の方法。

〔例 7〕

上記シーケンスの各トークンに対応する列エンコーディングを決定するステップと、
上記入力シーケンスの各トークンについてのスカラー注目スコアを含むベクトルを決定するステップと、

上記ベクトルをソフトマックス関数を使用して正規化するステップと、

30

上記入力表現を、対応する正規化されたスコアにより重み付けされた上記列エンコーディングの和として決定するステップと、

をさらに含む例 5 に記載の方法。

〔例 8〕

上記複数のモデルを勾配降下を使用して訓練して、上記複数のモデルの各々の結果に基づいて損失を表す目的関数を最小化するステップ

をさらに含む例 1 に記載の方法。

〔例 9〕

上記 1 つ以上の入力表現を生成するステップは、入力表現 κ^{agg} を、

トークンのシーケンス内の各 t 番目のトークンについて、スカラー注目スコア α^{in}
 $t = W^{inp} h^{enc}_t$ を計算し、 h^{enc}_t は、上記入力シーケンス内の t 番目のワードに対応するエンコードの状態であり、

40

上記スコアのベクトル $\alpha^{in} = [\alpha^{inp}_1, \alpha^{inp}_2, \dots]$ を正規化して、
上記トークンのシーケンス内の上記トークンにわたる分布を生成し、

上記入力表現 κ^{agg} を、

【数 7】

$$\kappa^{agg} = \sum_t \beta_t^{inp} h_t^{enc}$$

として取得し、 $\beta^{agg} = \text{softmax}(\alpha^{agg})$ 及び $\alpha^{agg} = W^{agg} \tan h$

50

($V^a g g_k^a g g + b^a g g$) + $c^a g g$ である

ことにより計算するステップを含む、例 1 乃至 8 のうちいずれか 1 項に記載の方法。

〔例 10〕

1 つ以上の列名を有する上記 S Q L クエリの $select$ 句を定式化する (formulates) ポインタネットワークで長短期記憶 (LSTM) を使用するステップであり、

列表現のリスト及び質問表現を所与として、質問に最もマッチする列を選択することを含み、上記列表現のリストは、各列名を LSTM でエンコードすることにより取得され、特定の列 j の表現 e_j^c は、

【数 8】

$$h_{j,t}^c = \text{LSTM}(\text{emb}(x_{j,t}^c), h_{j,t-1}^c) \quad e_j^c = h_{j,T_j}^c$$

10

により与えられる、ステップ、

をさらに含む例 1 乃至 9 のうちいずれか 1 項に記載の方法。

〔例 11〕

強化学習を使用して、データベースに対して上記生成された S Q L クエリを実行することにより上記 S Q L クエリの $where$ 条件を定式化して、

【数 9】

$$R(q(y), q_g) = \begin{cases} -2, & q(y) \text{ が有効な S Q L クエリでない場合} \\ -1, & q(y) \text{ が有効な S Q L クエリであり、不正確な結果に対し実行する場合} \\ +1, & q(y) \text{ が有効な S Q L クエリであり、正確な結果に対し実行する場合} \end{cases}$$

20

として定義される報酬 $R(q(y), q_g)$ を取得するステップであり、 $q(y)$ は、上記モデルにより生成されたクエリを表し、 q_g は、上記入力自然言語クエリに対応するグラウンドトゥースクエリを表す、ステップ

をさらに含む例 1 乃至 10 のうちいずれか 1 項に記載の方法。

〔例 12〕

自然言語要求を処理して S Q L クエリの部分を生成するステップであり、

1 つ以上の列名を有する上記 S Q L クエリの $select$ 句を定式化するポインタネットワークで長短期記憶 (LSTM) を使用することと、

強化学習により訓練された拡大ポインタデコーダ (augmented pointer decoder) を使用して上記 S Q L クエリの $where$ 条件を定式化するステップと、

を含む方法。

〔例 13〕

多層パーセプトロンを使用して、指定された条件下で選択された列に適用可能な上記 S Q L クエリの集約演算を決定するステップ、

をさらに含む例 12 に記載の方法。

〔例 14〕

1 つ以上の列名を有する上記 S Q L クエリの $select$ 句を定式化するポインタネットワークで長短期記憶 (LSTM) を使用することは、

列表現のリスト及び質問表現を所与として、質問に最もマッチする列を選択することを含む、例 12 又は 13 に記載の方法。

〔例 15〕

上記列表現のリストは、各列名を LSTM でエンコードすることにより取得され、特定の列 j の表現 e_j^c は、

【数 10】

$$h_{j,t}^c = \text{LSTM}(\text{emb}(x_{j,t}^c), h_{j,t-1}^c) \quad e_j^c = h_{j,T_j}^c$$

により与えられる、例 14 に記載の方法。

50

〔例 16〕

上記 1 つ以上の入力表現を生成することは、入力表現 κ^{agg} を、

トークンのシーケンス内の各 t 番目のトークンについて、スカラー注目スコア $\alpha^{in}_t = W^{inp} h^{enc}_t$ を計算し、 h^{enc}_t は、上記入力シーケンス内の t 番目のワードに対応するエンコーダの状態であり、

上記スコアのベクトル $\alpha^{in} = [\alpha^{inp}_1, \alpha^{inp}_2, \dots]$ を正規化して、上記トークンのシーケンス内の上記トークンにわたる分布を生成し、

上記入力表現 κ^{agg} を、

【数 11】

$$\kappa^{agg} = \sum_t \beta_t^{inp} h_t^{enc}$$

10

として取得し、 $\beta^{agg} = \text{softmax}(\alpha^{agg})$ 及び $\alpha^{agg} = W^{agg} \tanh(V^{agg} \kappa^{agg} + b^{agg}) + c^{agg}$ である

ことにより計算することを含む、例 14 乃至 15 のうちいずれか 1 項に記載の方法。

〔例 17〕

強化学習を使用して、データベースに対して上記生成された SQL クエリを実行することにより上記 SQL クエリの where 条件を定式化して、

【数 12】

$$R(q(y), q_g) = \begin{cases} -2, & q(y) \text{ が有効な SQL クエリでない場合} \\ -1, & q(y) \text{ が有効な SQL クエリであり、不正確な結果に対し実行する場合} \\ +1, & q(y) \text{ が有効な SQL クエリであり、正確な結果に対し実行する場合} \end{cases}$$

20

として定義される報酬 $R(q(y), q_g)$ を取得するステップであり、 $q(y)$ は、上記モデルにより生成されたクエリを表し、 q_g は、上記入力自然言語クエリに対応するグラウンドトゥースクエリを表す、ステップ

をさらに含む例 14 乃至 16 のうちいずれか 1 項に記載の方法。

〔例 18〕

1 つ以上のプロセッサにより実行されたときに上記 1 つ以上のプロセッサに例 1 乃至 17 のうちいずれか 1 項に記載のステップを実行させるコンピュータ実行可能コードを含む非一時的コンピュータ読取可能記憶媒体。

30

〔例 19〕

リカレントニューラルネットワーク (RNN) アーキテクチャを生成するコンピュータシステムであって、

1 つ以上のコンピュータプロセッサと、

1 つ以上のコンピュータにより実行されたときに上記 1 つ以上のコンピュータに請求項 18 に記載の方法のステップを実行させるコンピュータ実行可能コードを含む非一時的コンピュータ読取可能記憶媒体と、

を含むコンピュータシステム。

〔例 20〕

40

1 つ以上のコンピュータプロセッサと少なくとも 1 つの非一時的記憶媒体とを含むコンピュータシステムであって、

データベーススキーマを使用して記憶されたデータに基づく入力自然言語クエリを受信する手段と、

上記入力自然言語クエリのターム、

上記データベーススキーマの列のセット、及び、

データベースクエリ言語の語彙

を含む複数のタームからトークンのシーケンスを生成する手段と、

1 つ以上の入力表現を生成する手段であり、各入力表現は、上記トークンのシーケンスをエンコードすることにより取得される、手段と、

50

複数の機械学習に基づくモデルにアクセスする手段であり、各モデルは、上記入力自然言語クエリに対応するデータベースクエリの部分を予測するように構成される、手段と、

上記複数のモデルの各々について、上記モデルを入力表現に基づいて実行して上記データベースクエリの部分を生成する手段と、

上記データベースクエリの上記生成された部分を組み合わせて上記データベースクエリを取得する手段と、

上記データベースクエリを実行して結果セットを取得する手段と、

をさらに含むコンピュータシステム。

〔例 2 1〕

1 つ以上のコンピュータプロセッサと少なくとも 1 つの非一時的記憶媒体とを含むコンピュータシステムであって、

自然言語要求を処理して SQL クエリの部分を生成する手段であり、

1 つ以上の列名を有する上記 SQL クエリの select 句を定式化するポインタネットワークで長短期記憶 (LSTM) を使用する手段と、

強化学習により訓練された拡大ポインタデコーダを使用して上記 SQL クエリの where 条件を定式化する手段と、

を含む、手段

をさらに含むコンピュータシステム。

10

【0058】

20

〔代替的な実施形態〕

開示された実施形態は、関係データベースに基づいており、SQL を使用して例示されているが、開示された技術は、他のタイプのデータベース、例えば、オブジェクトベースのデータベース、オブジェクト関係データベースなどに適用可能である。開示された技術は、特定タイプのデータベースに使用されるデータベースクエリ言語が結果列、集約句、又は条件句に相当する機能をサポートする場合、適用可能である。例えば、データベースクエリ言語が条件句をサポートする場合、条件句予測器を使用して、入力自然言語クエリに基づいて出力データベースクエリの条件句を予測することができる。

【0059】

本発明の図及び説明は、明確さの目的で、典型的な分散システムに見られる他の多くの要素を排除すると同時に、本発明の明確な理解に関連する要素を例示するよう簡略化されていることを理解されたい。当業者は、他の要素及び／又はステップが実施形態を実現する際に望ましく、かつ／あるいは必要とされることを認識し得る。しかしながら、こうした要素及びステップは当該分野において良く知られているため、かつそれらは実施形態のより良い理解を促進しないため、こうした要素及びステップの議論は本明細書で提供されない。本明細書の開示は、当業者に知られるこうした要素及び方法に対するすべてのこうした変形及び修正に向けられる。

30

【0060】

上記説明のいくつかの部分は、実施形態を、情報に対する演算のアルゴリズム及びシンボル表現の観点から説明する。これらのアルゴリズムの説明及び表現は、データ処理分野の当業者により一般的に使用され、その作用の実体を他の当業者に効果的に伝達する。これらの動作は、機能的、計算的、又は論理的に説明されているが、コンピュータプログラム又は同等の電気回路、マイクロコードなどにより実現されるものと理解される。さらに、これらの動作の配置を一般性を損なうことなくモジュールとして参照することは、時に便利であることも証明されている。説明した動作及びその関連モジュールは、ソフトウェア、ファームウェア、ハードウェア、又はこれらの任意の組み合わせで具体化されてもよい。

40

【0061】

本明細書で用いられるとき、「1 つの実施形態」又は「一実施形態」へのいずれの参照も、実施形態に関連して説明された特定の要素、特徴、構造、又は特性が少なくとも 1 つ

50

の実施形態に含まれることを意味する。明細書中の様々な箇所におけるフレーズ「1つの実施形態において」の出現は、必ずしもすべてが同じ実施形態を参照しているわけではない。

【0062】

いくつかの実施形態は、表現「結合された」及び「接続された」をその派生形と共に使用して説明され得る。これらの用語は、互いに同義語として意図されてはいないことを理解されたい。例えば、いくつかの実施形態は、用語「接続された」を使用して説明されて、2つ以上の要素が互いに直接物理的又は電氣的に接触していることを示し得る。別の例において、いくつかの実施形態は、用語「結合された」を使用して説明されて、2つ以上の要素が直接物理的又は電氣的に接触していることを示し得る。しかしながら、用語「結合された」は、2つ以上の要素が互いに直接接触していないが、なお依然として互いに協調又は対話することをさらに意味し得る。実施形態は、この文脈において限定されない。

10

【0063】

本明細書で用いられるとき、用語「含む」、「含んでいる」、「含める」、「含めている」、「有する」、「有している」、又はこれらの任意の他の変形は、非排他的な包含をカバーすることが意図される。例えば、要素のリストを含む処理、方法、物品、又は装置は、必ずしもこれらの要素のみに限定されず、明確に列挙されていないか又はこうした処理、方法、物品、若しくは装置に固有の他の要素を含んでもよい。さらに、逆のことを明確に示されない限り、「又は」は、排他的ORでなく包括的ORを参照する。例えば、条件A又はBは、Aが真であり（又は存在し）Bが偽である（又は存在しない）と、Aが偽であり（又は存在せず）Bが真である（又は存在する）と、A及びBの双方が真である（又は存在する）と、のうちいずれか1つにより満たされる。

20

【0064】

さらに、一の（“a”又は“an”）の使用は、本明細書における実施形態の要素及びコンポーネントを説明するために採用されている。これは単に、簡便さのため、及び本発明の一般的な意味を与えるために行われる。本説明は、1つ又は少なくとも1つを含むように読まれるべきであり、その他を意図されることが明らかでない限り、単数形は複数形も含む。

【0065】

本開示を読むと、当業者は、本明細書で開示された原理を通して、歪み領域を使用してチャートを表示するシステム及び処理のためのなおさらなる代替的な構造及び機能設計を理解するであろう。ゆえに、特定の実施形態及び適用が例示され説明されたが、開示された実施形態は、本明細書に開示された正確な構造及びコンポーネントに限定されないことが理解されるべきである。当業者に明らかであろう様々な修正、変更及び変形が、別記の特許請求の範囲に定義された主旨及び範囲から逸脱することなく、本明細書に開示された方法及び装置の配置、動作、及び詳細において行われ得る。

30

【0066】

[付録]

以下の「Seq2SQL: Generating Structured Queries from Natural Language using Reinforcement Learning」と題された文献は、本付録の一部であり、ゆえに本出願の開示の一部である「Seq2SQL」を開示しており、「Seq2SQL」、自然言語の質問に対応するSQLクエリに翻訳するためのディープニューラルネットワークを開示している。Seq2SQLは、SQLクエリの構造を活用して、生成されたクエリの出力空間を有意に削減する。Seq2SQLは、SQLクエリの構造を活用する3つのコンポーネントを含み、生成されたクエリの出力空間を大幅に削減する。詳細には、Seq2SQLは、ポイントデコード及びポリシー勾配を使用してクエリの条件を生成し、これは、順序づけられていない性質に起因して交差エントロピー損失を使用した最適化に適さないことを我々は示す。Seq2SQLは、交差エントロピー損失とデータベースに対するライブのクエリ実行からの強化学習報酬とを組み合わせた混合された目的を使用して訓練される。これらの特性は、モデルがクエリ生成において一層向上した結果を達成することを

40

50

可能にする。

WO 2018/213530

23

PCT/US2018/033099

Seq2SQL: Generating Structured Queries from Natural Language using Reinforcement Learning

Anonymous Author(s)
Affiliation
Address
email

Abstract

A significant amount of the world's knowledge is stored in relational databases. However, the ability for users to retrieve facts from a database is limited due to a lack of understanding of query languages such as SQL. We propose Seq2SQL, a deep neural network for translating natural language questions to corresponding SQL queries. Our model leverages the structure of SQL queries to significantly reduce the output space of generated queries. Moreover we apply policy gradient reinforcement learning using in-the-loop query execution to learn a policy for generating the conditions of the SQL query. Elements of queries are unordered and not suitable for optimization via cross entropy loss. In addition, we will publish WikiSQL, a dataset of 37,000 hand-annotated examples of questions and SQL queries distributed across 9,500 tables from Wikipedia. This dataset is required to train our model and is an order of magnitude larger than comparable datasets. By applying policy gradient methods using rewards from in-the-loop query execution on WikiSQL, we show that Seq2SQL outperforms attentional sequence-to-sequence models, improving execution accuracy from 21.8% to 52.2% and logical form accuracy from 20.4% to 46.6%.

1 Introduction

A vast amount of today's information is stored in relational databases, which provide the foundation of crucial applications such as medical records [Hillestad et al., 2005], financial markets [Beck et al., 2000], and customer relations management [Ng et al., 2009]. However, accessing information in relational databases requires an understanding of query languages such as SQL, which, while powerful, is difficult to master. A prominent research area at the intersection of natural language processing and human-computer interactions is the study of natural language interfaces (NLI) [Andrioutsopoulou et al., 1993], which seek to provide means for humans to interact with computers through the use of natural language. We investigate one particular aspect of natural language interfaces applied to relational databases: translating natural language questions to SQL queries.

Our main contributions in this work are two-fold. First, we introduce Seq2SQL, a deep neural network for translating natural language questions to corresponding SQL queries. Seq2SQL, shown in Figure 1, consists of three components that leverage the structure of a SQL query to greatly reduce the output space of generated queries. In particular, it uses a pointer decoder and policy gradient to generate the conditions of the query, which we show are unsuitable for optimization using cross entropy loss due to their unordered nature. Seq2SQL is trained using a mixed objective combining cross entropy losses and reinforcement learning rewards from live query execution on a database. These characteristics allow the model to achieve much improved results on query generation.

Second, we release WikiSQL, a corpus of 37,000 hand-annotated instances of natural language questions, SQL queries, and SQL tables extracted from 9,500 HTML tables from Wikipedia. WikiSQL

Submitted to 34th Conference on Neural Information Processing Systems (NIPS 2017). Do not distribute.

WO 2018/213530

24

PCT/US2018/033099

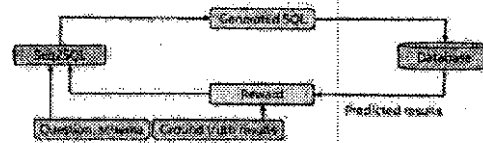


Figure 1: Seq2SQL takes as input a question and list of columns of a table. It generates the corresponding SQL query, which, during training, is executed against a database. The result of the execution is utilized as the reward to train the reinforcement learning algorithm.

is an order of magnitude larger than comparable datasets which also provide logical forms along with natural language utterances. We release the tables used in WikisQL both in raw JSON format as well as in the form of a SQL database. Along with WikisQL, we also release a query execution engine for the database, which we use for in-the-loop query execution during policy learning. On WikisQL, we show that Seq2SQL significantly outperforms an attentional sequence-to-sequence baseline, which obtains 20.3% execution accuracy, as well as a pointer network baseline, which obtains 43.2% execution accuracy. By leveraging the inherent structure of SQL queries and applying policy gradient methods using the reward signals from live query execution during training, Seq2SQL achieves a marked improvement on WikisQL, obtaining 52.2% execution accuracy.

2 Related Work

Learning logical forms from natural language utterances. In semantic parsing for question answering (QA), natural language questions are parsed into logical forms, which are then executed on a knowledge graph (Zelle and Mooney, 1996; Wong and Mooney, 2007; Zettlemoyer and Collins, 2005, 2007). Other work in semantic parsing has focused on learning parsers without relying on annotated logical forms by leveraging conversational logs (Arni and Zettlemoyer, 2011), demonstrations (Arni and Zettlemoyer, 2013), distant supervision (Cai and Yates, 2013; Reddy et al., 2014), and question-answer pairs (Liang et al., 2011). Recently, Pasupat and Liang (2015) introduced a framing parser to address the large amount of variation in utterances and table schema for semantic parsing on web tables. Our approach is different from these semantic parsing techniques in that it does not rely on a high quality grammar and lexicon, and instead relies on representation learning.

Semantic parsing datasets. Previous semantic parsing systems were typically designed to answer complex and compositional questions over close-domain, fixed-schema datasets such as Geo-Query (Yang and Mooney, 2001) and ATIS (Price, 1990). Researchers have also investigated QA over subsets of large-scale knowledge graphs such as DBPedia (Starc and Mladenic, 2017) and Freebase (Cai and Yates, 2013; Berant et al., 2013). Compared to these datasets, WikisQL is the largest by a significant margin - it is an order of magnitude larger in the number of examples and/or the number of domains/table schema. The dataset "Overnight" (Wang et al., 2013) uses a similar crowd-sourcing process to build a dataset of (natural language question, logical form) pairs, but it is comprised of only 8 domains. WikitablesQuestions (Pasupat and Liang, 2015), which like WikisQL, is also a collection of question and answers over a large quantity of tables extracted from Wikipedia, achieves a balance between the breadth of the problem domain and the depth of question complexity. However WikitablesQuestions does not provide logical forms whereas WikisQL does. In addition, WikitablesQuestions is focused on the task of QA over potentially noisy web tables, whereas WikisQL is focused on generating SQL queries for questions over relational database tables, with the intent of building natural language interfaces for real-world, production databases.

Representation learning for sequence generation. Our baseline model is based on the attentional sequence-to-sequence model (Seq2Seq) proposed by (Babdan et al., 2015). Seq2Seq has led to significant progress in a variety of natural language processing tasks such as machine translation (Sutskever et al., 2014; Babdan et al., 2015), summarization (Nallapati et al., 2016; Paulus et al., 2017), and semantic parsing (Dew and Lapata, 2016). Seq2SQL is inspired by pointer networks (Vinyals et al., 2015), which, instead of generating words from a fixed vocabulary like attentional sequence-to-sequence models, generates by selecting words from the input sequence. This class of models has been successfully applied to tasks such as language modeling (Merity et al., 2017), summarization (Chen et al., 2016), combinatorial optimization (Bello et al., 2017), and question answering (Seo et al., 2017; Xiong et al., 2017). Neural methods have also been applied to semantic

WO 2018/213530

25

PCT/US2018/033099

Rank #	City/Team	Player	Position	League
27	San Antonio Spurs	Tim Duncan	PF	NBA
28	Los Angeles Lakers	Kobe Bryant	SG	NBA
29	Los Angeles Lakers	LeBron James	PF	NBA
30	Los Angeles Lakers	Shaquille O'Neal	C	NBA

Question: How many players are there in the NBA?

SQL: `SELECT COUNT(*) FROM Table WHERE League = 'NBA';`

Result: 30

Figure 2: An example in WikiSQL. The inputs consist of a table and question. The outputs consist of a ground truth SQL query and the corresponding result from execution.

22 parsing (Dong and Lapata, 2016; Neelakantan et al., 2017). However, unlike (Dong and Lapata, 2016),
 23 we use pointer based generation instead of Seq2Seq generation and show that the former approach
 24 works significantly better, especially for unseen words and column names. Unlike (Neelakantan et al.,
 25 2017), our model does not require access to the table content during inference. In contrast to both
 26 methods, we train our model using policy gradient, which again significantly improves performance.

27 **Natural language interface for databases.** The practical applications of WikiSQL have much
 28 to do with Natural Language Interfaces (NLI). One of the prominent works in this area is PRE-
 29 CISE (Pyrescu et al., 2003), which translates questions to SQL queries and identifies questions that it
 30 is not confident about. Giordani and Muschik (2012) proposed another system to translate questions
 31 to SQL, in which the model first generates candidate queries from a grammar, then ranks the queries
 32 using tree kernels. Both of these approaches rely on a high quality grammar and are likely unsuitable
 33 for tasks such as WikiSQL, which requires generalization to new schema. Ayer et al. (2017) also
 34 proposed a system to translate in SQL, but with a Seq2Seq model that is further improved with human
 35 feedback. Compared to this model, we show that Seq2SQL substantially outperforms Seq2Seq and
 36 uses reinforcement learning instead of human feedback during training.

27 3 Model

38 The Seq2SQL model generates a SQL query from a natural language question and table schema.
 39 One powerful baseline model is the Seq2Seq model utilized for machine translation (Bahdanau et al.,
 40 2015). However, the output space of the standard softmax in a Seq2Seq model is unnecessarily large
 41 for this task. In particular, we note that the problem can be dramatically simplified by limiting the
 42 output space of the generated sequence to the union of the table schema, the question inference,
 43 and SQL key words. The resulting model is similar to a pointer network (Vinyals et al., 2015) with
 44 augmented inputs. We first show this augmented pointer network model, then address its limitations,
 45 particularly with respect to generating unordered query conditions, in our description of Seq2SQL.

28 3.1 Augmented Pointer Network

46 The augmented pointer network generates the SQL query token-by-token by selecting from an input
 47 sequence. In our case, the input sequence is the concatenation of the column names, required for the
 48 selection column and the condition columns of the query, the question, required for the conditions
 49 of the query, and the limited vocabulary of the SQL language such as SELECT, COUNT etc. In the
 50 example shown in Figure 2, the column name tokens consist of "Pick", "P", "CFL", "Team" etc.; the
 51 question tokens consist of "How", "many", "CH", "teams" etc.; the SQL tokens consist of SELECT,
 52 WHERE, COUNT, MIN, MAX etc. With this augmented input sequence, the pointer network is able to
 53 fully produce the SQL query by selecting exclusively from the input.

54 Suppose we are given a list of N table columns and a question such as in Figure 2, and asked to
 55 produce the corresponding SQL query. Let $x_j^c = [x_j^c, x_j^q, x_j^s]$ denote the sequence of words in
 56 the name of the j th column, where x_j^c represents the i th word in the j th column and T_j represents
 57 the total number of words in the j th column. Similarly, let x^q and x^s respectively denote the sequence
 58 of words in the question and the set of all unique words in the SQL vocabulary.

59 We define the input sequence x as the concatenation of all the column names, the question, and the
 60 SQL vocabulary:

$$x = [x^c, x^q, x^s] \quad (1)$$

For simplicity, we define table schema as the names of the columns in the table.

WO 2018/213530

26

PCT/US2018/033099

where $\| \cdot \|$ denotes the concatenation between the sequences u and v and we add special sentinel tokens between all three sequences to demarcate the boundaries.

The network first encodes x using a two-layer, bidirectional Long Short-Term Memory network [Hochreiter and Schmidhuber, 1997]. The input to the encoder are the embeddings corresponding to words in the input sequence. We denote the output of the encoder by h_i where h_i is the state of the encoder corresponding to the i th word in the input sequence. For brevity, we do not write out the LSTM equations, which are described by Hochreiter and Schmidhuber [1997]. We then apply a pointer network similar to that proposed by Vinyals et al. [2015] to the input encodings h_i .

The decoder network uses a two-layer, unidirectional LSTM. During each decoder step s , the decoder LSTM takes as input q_{s-1} , the query token generated during the previous decoding step, and outputs the state g_s . Next, the decoder produces a scalar attention score $\alpha_{i,s}^{att}$ for each position i of the input sequence x :

$$\alpha_{i,s}^{att} = W^{att} \tanh(U^{att} g_s + V^{att} h_i) \quad (2)$$

The input token with the highest score is then chosen by the decoder to be the next token of the SQL query: $h_s = \arg\max(\alpha_{i,s}^{att})$.

3.2 Seq2SQL

While the augmented pointer model is able to solve the SQL generation problem, it does not leverage the structure inherent in SQL queries. Normally, a SQL query such as that shown in Figure 3 consists of three components. The first component is the aggregation operator, in this case COUNT, which produces a summary of the rows selected by the SQL query. Alternatively the query may request no summary statistics, in which case an aggregation operator is not provided. The second component is the selection column(s), in this case Engine, which identifies the column(s) that are to be included in the returned results. The third component is the WHERE clause of the query, in this case WHERE Driver = Val Musetti, which contains conditions by which to filter the rows. Here, we want only rows in which the driver is "Val Musetti".

To leverage the structure present in SQL queries, we introduce Seq2SQL. Seq2SQL, as shown in Figure 3, is composed of three parts that correspond to the aggregation operator, the selection column, and the WHERE clause. First, the network classifies an aggregation operation for the query, with the addition of a null operation that corresponds to no aggregation. Next, the network points to a column in the input table corresponding to the SELECT column. Finally, the network generates the conditions for the query using a pointer network. The first two components are supervised using cross entropy loss, whereas the third generation component is trained using policy gradient reinforcement learning in order to address the unordered nature of query conditions (we explain this in detail in Section 3.2). Similar to moving from Seq2Seq to the augmented pointer model, utilizing the structure of a SQL query allows Seq2SQL to further reduce the output space of SQL queries. We show later that this leads to a significantly higher performance compared to both the Seq2Seq model and the pointer model.

Aggregation Operation

The aggregation operation depends on the question. For the example shown in Figure 3, we intuit that the correct operator is COUNT because the question asks for "How many".

To compute the aggregation operation, we first compute an input representation x^{enc} . Similar to Equation 2, we first compute the scalar attention score, $\alpha_{i,s}^{att} = W^{att} g_s + V^{att} h_i$, for each i th token in the input sequence. The vector of all scores, $\alpha^{att} = [\alpha_1^{att}, \alpha_2^{att}, \dots]$, is then normalized to produce a distribution over all the tokens in the input sequence x : $\beta^{att} = \text{softmax}(\alpha^{att})$. The input representation, x^{enc} , is the sum over the column encodings weighted by the corresponding

WO 2018/213530

27

PCT/US2018/033099

175 normalized attention weight:

$$176 \quad \alpha^{agg} = \frac{\sum_j \beta_j^{agg} v_j^{agg}}{\sum_j \beta_j^{agg}} \quad (3)$$

177 Let α^{agg} denote the scores over the aggregation operations such as COUNT, MIN, MAX, and the no-
178 aggregation operation NULL. We compute α^{agg} by applying a multi-layer perceptron to the input
179 representation v^{agg} :

$$180 \quad \alpha^{agg} = W^{agg} \text{tanh}(V^{agg} v^{agg} + b^{agg}) + e^{agg} \quad (4)$$

181 The distribution over the set of possible aggregation operations, $\beta^{agg} = \text{softmax}(\alpha^{agg})$, is then
182 obtained by applying the softmax function. We use cross entropy loss L^{agg} for the aggregation
183 operation.

184 **SELECT Column**

185 The selection column depends on the table columns as well as the question. Namely, for the example
186 in Figure 3, we see from the "How many engine types" part of the question that the query is asking
187 for the "Engine" column. SELECT column prediction is then a matching problem, solvable using a
188 pointer: given the list of column representation and a question representation, we select the column
189 that best matches the question.

190 In order to produce the representations for the columns, we first uncode each column name with a
191 LSTM. The representation of a particular column j , e_j^c , is given by:

$$192 \quad h_{j,t}^c = \text{LSTM}(\text{emb}(e_{j,t}^c), h_{j,t-1}^c) \quad e_j^c = h_{j,T}^c \quad (5)$$

193 Here, $h_{j,t}^c$ denotes the t th encoder state of the j th column. e_j^c , column j 's representation, is then taken
194 to be the last encoder state.

195 To construct a representation for the question, we compute another input representation v^{sel} using the
196 same architecture as for v^{agg} (Equation 3) but with unified weights. Finally, we apply a multi-layer
197 perceptron over the column representations for the pointer estimation, conditioned on the input
198 representation, to compute the score for each column j :

$$199 \quad \alpha_j^{sel} = W^{sel} \text{tanh}(V^{sel} v^{sel} + V^{sel} e_j^c) \quad (6)$$

200 We normalize the scores with a softmax function to produce a distribution over the possible SELECT
201 columns $\beta^{sel} = \text{softmax}(\alpha^{sel})$. For the example shown in Figure 3, the distribution would be over
202 the columns "Entrant", "Constructor", "Chassis", "Engine", "No", and the ground truth SELECT
203 column "Driver". We train the SELECT network using cross entropy loss L^{sel} .

204 **WHERE Clause**

205 The WHERE clause can be trained using a pointer decoder similar to that described in Section 3.1.
206 However, there is a crucial limitation in using the cross entropy loss to optimize the network: the
207 WHERE conditions of a query can be arbitrarily swapped and the query still yield the same result.
208 Suppose we have the question "what are the names of men who are older than 18" and the queries:
209 SELECT name FROM insurance WHERE age > 18 AND gender = "male" and SELECT name
210 FROM insurance WHERE gender = "male" AND age > 18. Both queries obtain the correct
211 execution result despite not having exact string match. Suppose the former query is provided as the
212 ground truth, using cross entropy loss to supervise the generation would then wrongly penalize the
213 latter query. To address this problem, we apply reinforcement learning to learn a policy to directly
214 optimize the expected "correctness" of the execution result (Equation 7).

215 Instead of teacher forcing at each step, we sample from the output distribution to obtain the next
216 token. At the end of the generation procedure, we execute the generated SQL query against the
217 database to obtain a reward. Let $q(y)$ denote the query generated by the model and q_g denote the
218 ground truth query corresponding to the question. The reward $R(q(y), q_g)$ is defined as:

$$219 \quad R(q(y), q_g) = \begin{cases} -2, & \text{if } q(y) \text{ is not a valid SQL query} \\ -1, & \text{if } q(y) \text{ is a valid SQL query and executes to an incorrect result} \\ +1, & \text{if } q(y) \text{ is a valid SQL query and executes to the correct result} \end{cases} \quad (7)$$

WO 2018/213530

28

PCT/US2018/033099

Let $y = [y^1, y^2, \dots, y^T]$ denote the sequence of generated tokens in the WHERE clause of a SQL query $q(y)$ that resulted in an execution reward of $R(q(y), q_a)$. The loss, $L^{wre} = -\mathbb{E}_y[R(q(y), q_a)]$, is the negative expected reward over possible WHERE clauses.

As shown by Schulman et al. [2015], the act of sampling during the forward pass of the network results in the corresponding injection of a synthetic gradient, which is a function of the reward, during the backward pass. Specifically, let $P(y^t)$ denote the probability of choosing token y^t during time step t ; the corresponding synthetic gradient is $R(q(y), q_a) \log P(y^t)$.

Mixed Objective Function

The model is trained using gradient descent to minimize the objective function $J^{tot} = L^{sel} + L^{wre} + L^{agg}$. The total gradient during the backward pass is then the equally weighted sum of the gradients from the cross entropy loss in predicting the SELECT column, the cross entropy loss in predicting the aggregation operation, and the expected reward for policy learning as described above.

4 WIKISQL

WikisQL is a collection of natural language utterances and corresponding SQL queries and SQL tables. To our knowledge, it is the largest corpus available of its kind. A single example in WikisQL, such as that shown in Figure 2, contains a SQL table, a SQL query, as well as the natural language utterance corresponding to the SQL query. In this work, we do not use the content of the SQL table, but only its schema (e.g. column names). This means that systems trained on WikisQL do not need to access the end user's database contents during inference.

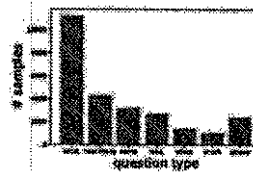


Figure 4: Distribution of question types in WikisQL.

Table 1 shows how WikisQL compares to related datasets. Namely, we note that WikisQL is the largest hand-annotated semantic parsing dataset to date. It is an order of magnitude larger than other datasets that have logical forms, either in terms of the number of examples or the number of table schemas. Moreover, the queries in WikisQL span over a large number of table schemas and cover a large diversity of data. The large quantity of schemas presents a challenge compared to other datasets: the model must be able to not only generalize to new queries, but to new schemas as well. Finally, WikisQL contains challenging, realistic data extracted from the web. This is evident in the distributions of the number of columns, the lengths of questions, and the length of queries, which are respectively shown in Figure 3. Another indicator of the variety of questions in the dataset is the distribution of question types, which is shown in Figure 4.

WikisQL is collected using crowd-sourcing on Amazon Mechanical Turk (AMT) in two phases. In the first phase, a worker is asked to paraphrase a generated question for a table. The generated question itself is formed using a template, filled using a randomly generated SQL query. Due to the large variation in HTML tables, we ensure the validity and complexity of the tables by keeping only those that are legitimate database tables and are sufficiently large in the number of rows and columns. In the second phase, two other workers are asked to verify that the paraphrase has the same meaning as the generated question. We discard paraphrases that do not show enough variation, as measured by the character edit distance from the generated question, as well as those that are deemed incorrect by both workers during the verification phase. More details on how WikisQL was collected is shown in Section A of the Appendix. We make available examples of the interfaces used during the paraphrase phase² and during the verification phase³. The HTML pages of the interfaces are also included in the supplementary materials.

The tables and their corresponding paraphrases and SQL queries are randomly slotted into train, dev, and test splits, such that each table is only present in exactly one split. In addition to the raw content

²Example of the paraphrasing interface used for WikisQL: <https://tablemap.herokuapp.com/example/3652872-26>

³Example of the verification interface used for WikisQL: [https://paraphrase.herokuapp.com/example/3652872-26?question=SELECT%20COUNT\(*\)%20FROM%20TABLE1%20WHERE%20ID=1](https://paraphrase.herokuapp.com/example/3652872-26?question=SELECT%20COUNT(*)%20FROM%20TABLE1%20WHERE%20ID=1)

WO 2018/213530

29

PCT/US2018/033099

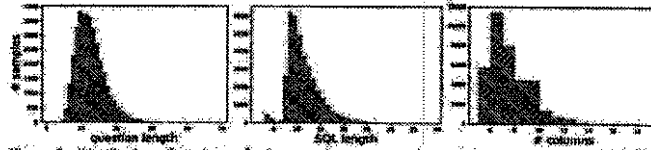


Figure 5: Distribution of numbers of columns, tokens per question, and tokens per query in WikisQL.

of tables, queries, results, and natural utterances, we also release a corresponding SQL database as well as an execution engine for queries system.

4.1 Evaluation

One natural way to evaluate WikisQL is to use the execution accuracy metric. Let N denote the total number of examples in the dataset and N_{acc} denote the total number of queries that, when executed, result in the correct result. The execution accuracy is defined as $\text{Acc}_{\text{ex}} = \frac{N_{\text{acc}}}{N}$.

Another metric is the logical form accuracy. Let N_{f} denote the total number of queries that exactly match with the ground truth query used to collect the paraphrase. The logical form accuracy is defined as $\text{Acc}_{\text{f}} = \frac{N_{\text{f}}}{N}$.

There are benefits and drawbacks to both metrics. As we will show in Section 3.2, the logical form accuracy incorrectly penalizes queries that achieve the correct result but do not have exact string match with the ground truth query. Alternatively, it is possible to construct a SQL query that does not correspond to the natural language question but nevertheless obtain the same result. An example in which two different queries provide the same result is `SELECT COUNT(name) WHERE SSN = 123` and `SELECT COUNT(SSN) WHERE SSN = 123` and `SELECT COUNT(SSN) WHERE SSN = 123`.

Dataset	Size	LF	Schema
WikisQL	37908	yes	9506
Geoquery	890	yes	8
ATIS	5871	yes	141
Freebase917	917	yes	81
Overnight	26098	yes	8
WebQuestions	3810	no	2420
WikiTableQuestions	22033	no	2108

Table 1: Comparison between WikisQL and existing datasets. The datasets are GeoQuery890 (Tang and Mooney, 2001), ATIS (Price, 1990), Free917 (Cai and Yates, 2013), Overnight (Wang et al., 2015), WebQuestions (Berant et al., 2013), and WikiTableQuestions (Pasupat and Liang, 2015). "Size" denotes the total number of examples in the dataset. "LF" indicates whether the dataset has annotated logical forms. "Schema" denotes the number of domains or schemas in the dataset. The format of ATIS is presented as a plot within the text.

5 Experiments

The dataset is tokenized using Stanford CoreNLP. In particular, we used the normalized tokens for training and reverse them into their original gloss before outputting the query. This way, the queries are composed of tokens from the original gloss and are hence executable in the database.

We use pretrained word embeddings from GloVe (Pennington et al., 2014) and character n-gram embeddings (Hashimoto et al., 2016). Let w_x denote the GloVe embedding and c_x the character embedding for word x . Here, c_x is the mean of the embeddings of all the character n -grams in x . For a given word x , we define its word embedding $v_x = [w_x; c_x]$ as the concatenation of its GloVe embedding and its character embedding. For words that do not have pretrained embeddings, we assign the embeddings to the zero vector. We do not train embeddings and instead fix the pretrained embeddings.

All networks are trained for a maximum of 100 epochs with early stopping using the execution accuracy criteria on the dev split. We train the networks using ADAM (Kingma and Ba, 2014) and regularization using dropout (Srivastava et al., 2014). We implement all models using PyTorch⁵.

⁵<https://pytorch.org>

WO 2018/213530

30

PCT/US2018/033099

Model	Dev Acc _g	Dev Acc _{RL}	Test Acc _g	Test Acc _{RL}
Attentional Seq2Seq	20.5%	22.4%	20.4%	21.8%
Aug. Pointer Network	38.3%	43.3%	37.7%	43.2%
Seq2SQL (w/ RL)	46.2%	49.9%	44.4%	47.3%
Seq2SQL	48.3%	54.2%	46.6%	52.2%

Table 2: Performance on WikisQL. Both metrics are defined in Section 4.1. The Seq2SQL (w/ RL) model is identical in architecture when compared to Seq2SQL, except the WHERE clause is supervised via teacher forcing as opposed to reinforcement learning.

To train Seq2SQL, we first pretrain a version in which the WHERE clause is supervised via teacher forcing (e.g. the sampling procedure is not trained from scratch) and then continue training using reinforcement learning. In order to obtain the rewards described in Section 3.2, we use the query execution engine described in Section 4.

5.1 Result

We compare results against an attentional sequence-to-sequence baseline used by Reymt et al. (2017) for machine translation. This model is described in detail in Section B of the Appendix. We also present an ablation of Seq2SQL, in which the WHERE clause of the SQL query is trained using the cross entropy loss (e.g. teacher forcing) in the same fashion as the augmented pointer network. The performance of the three models are shown in Table 2.

We note that reducing the output space by utilizing the augmented pointer network significantly improves upon the attentional sequence-to-sequence model by 21.4%. Moreover, leveraging the structure of SQL queries leads to another significant improvement of 4.1%, as is shown by the performance of Seq2SQL without RL compared to the augmented pointer network. Finally, the full Seq2SQL model, trained using reinforcement learning based on rewards from live query executions on a database, leads to yet another significant performance of 4.9%.

5.2 Analysis

Limiting the output space lead to more accurate conditions. When we compare the outputs of the attentional sequence-to-sequence model with those of the augmented pointer model, we find that the latter generates much higher quality WHERE clause conditions. For example, for the question "in how many districts was a successor seated on march 4, 1830?", former generates SELECT COUNT District WHERE Date successor seated = seated march 4 whereas the latter generates SELECT COUNT District WHERE Date successor seated = seated march 4, 1830. Similarly, for the question "what's doug battaglia's pick number?", the former generates SELECT Pick WHERE Player = doug whereas the latter generates SELECT Pick WHERE Player = doug battaglia. A key reason for this is that conditions tend to be comprised of rare words (e.g. "1830" occurs very infrequently in the corpus). The Seq2Seq model is then inclined to produce a condition that contains frequently occurring tokens in the training corpus, such as "march 4" for date, or "doug" for player name. The pointer network does not face this problem, as the output distribution is exclusively constructed from the input sequences.

Incorporating structure reduces invalid queries. As expected, incorporating the simple selection and aggregation structure into the model dramatically reduces the percentage of invalid SQL queries generated from 39.8% to 7.1%. Generally, we find that a large quantity of invalid queries result from invalid column names. That is, the model generates a query that refers to columns that are not present in the table. This is particularly helpful when the names of columns contain many tokens, such as "Miles (km)", which is comprised of 4 tokens after tokenization. Introducing a specialized classifier for the aggregation also reduces the error rate due to having the incorrect aggregation operator. For example, Table 3 shows that adding the aggregation classifier substantially improves the recall for predicting the COUNT operator. For more queries produced by the different models, please see Section C of the Appendix.

Reinforcement learning generates higher quality WHERE clause that are at times ordered differently compared to ground truth. As expected, the model trained with policy learning ob-

WO 2018/213530

PCT/US2018/033099

31

thus correct results in which the order of the conditions is different than the order in the ground truth query. For example, for the question "in what district was the democratic candidate first elected in 1992?", the ground truth conditions are `WHERE Party = Democratic AND First elected = 1992`. More generally, we note that in most cases where Seq2SQL is correct and SeqSQL without policy learning is not, the latter produces an incorrect `WHERE` clause. For example, for the rather complex question "what is the race name of the 12th round treston, new jersey race where a.j. foyt had the pole position?", the ablation of SeqSQL trained without reinforcement learning generates `WHERE race = 12 and track = a.j. foyt AND pole position = a.j. foyt` whereas the policy gradient based Seq2SQL correctly generates the clause `WHERE race = 12 AND pole position = a.j. foyt`.

Difficult table schemas present a challenge. In numerous HTML tables, the column names do not accurately portray the column contents. For example, in one of the tables in the dev set, the number of voters from the Bronx is listed under a column title "the bronx". When asked the question "how many voters from the bronx voted for the socialist party", the model mistakenly generates the aggregation operator `COUNT` due to the phrase "how many". However, due to the way this table is organized, the number of voters are actually listed in the column and the `COUNT` operation is superfluous. Another difficult example is the question "in what year did a plymouth vehicle win on february 9?", because the model does not understand that "plymouth" refers to a manufacturer of vehicles.

Model	Precision	Recall	F1
Avg. Pointer	87.0%	47.8%	61.7%
Seq2SQL	85.9%	79.7%	82.7%

Table 3: Performance on the COUNT aggregation operator. The Seq2SQL model's inclusion of the COUNT operator is compared with the ground truth query's inclusion of the COUNT operator.

6 Conclusion

We proposed Seq2SQL, a deep neural network for translating natural language questions to corresponding SQL queries that leverages the structure of SQL queries to significantly reduce the output space of generated queries. As a part of Seq2SQL, we proposed using in-the-loop query execution to learn a policy for generating the conditions of the SQL query, which is unordered in nature and not suitable for optimization via cross entropy loss. Along with Seq2SQL, we introduced WikiSQL, a dataset of natural language questions and SQL queries that is an order of magnitude larger than comparable datasets. Finally, we showed that Seq2SQL significantly outperforms attentional sequence to sequence models on WikiSQL, improving execution accuracy from 21.8% to 52.2% and logical form accuracy from 20.4% to 46.6%. In future work, we plan on investigating more complex queries such as nested queries and more sophisticated reward functions.

References

1. Andrioutsopoulos, G. Ritchie, and R. Thrun. Natural language interfaces to databases - an introduction. 1995.
2. Y. Artzi and L. S. Zettlemoyer. Bootstrapping semantic parsers from conversations. In *EMNLP*, 2011.
3. Y. Artzi and L. S. Zettlemoyer. Weakly supervised learning of semantic parsers for mapping instructions to actions. *TACL*, 1:49-62, 2013.
4. D. Bahdanau, K. Cho, and Y. Bengio. Neural machine translation by jointly learning to align and translate. *ICLR*, 2015.
5. T. Beck, A. Demirgüç-Kunt, and R. Levine. A new database on the structure and development of the financial sector. *The World Bank Economic Review*, 14(3):597-603, 2000.
6. I. Belle, H. Pham, Q. V. Le, M. Norouzi, and S. Bengio. Neural combinatorial optimization with reinforcement learning. *ICRL*, 2017.
7. F. Berant, A. Chou, R. Frostig, and B. Lian. Semantic parsing on freebase from question-answer pairs. In *EMNLP*, 2013.

WO 2018/213530

32

PCT/US2018/033099

- 385 C. Bhagavatula, T. Nowak, and D. Downey. Methods for exploring and mining tables on wikipedia.
386 In *IDEA@KDD*, 2013.
- 387 Q. Cai and A. Yates. Large-scale semantic parsing via schema matching and lexicon extension. In
388 *ACL*, 2013.
- 389 L. Deng and M. Lapata. Language to logical form with neural attention. *ACL*, 2016.
- 390 A. Giordani and A. Moschitti. Translating questions to SQL queries with generative parsers discrimin-
391 atively reranked. In *COLING*, 2012.
- 392 J. Gu, Z. Lu, H. Li, and V. O. K. Li. Incorporating copying mechanism in sequence-to-sequence
393 learning. *ACL*, 2016.
- 394 K. Hashimoto, C. Xiong, Y. Tsurumaki, and R. Socher. A Joint Many-Task Model: Growing a Neural
395 Network for Multiple NLP Tasks. *arXiv, cs.CL*, 1611.01587, 2016.
- 396 R. Hillestad, J. Bigelow, A. Bower, P. Girosi, R. Meili, R. Saville, and B. Taylor. Can electronic
397 medical record systems transform health care? potential health benefits, savings, and costs. *Health*
398 *affairs*, 24(5):1103-1117, 2005.
- 399 S. Hochreiter and J. Schmidhuber. Long short-term memory. *Neural computation*, 1997.
- 400 S. Iyer, I. Konstas, A. Cheung, J. Krishnamurthy, and L. Zettlemoyer. Learning a neural semantic
401 parser from user feedback. In *ACL*, 2017.
- 402 D. P. Kingma and J. Ba. Adam: A method for stochastic optimization. *arXiv, abs/1412.0980*, 2014.
- 403 G. Klein, Y. Kim, Y. Deng, J. Senellart, and A. M. Rush. OpenNMT: Open-Source Toolkit for Neural
404 Machine Translation. *ArXiv, abs/1701.02810*.
- 405 P. Liang, M. I. Jordan, and D. Klein. Learning dependency-based compositional semantics. *Compu-*
406 *tational Linguistics*, 39:389-446, 2011.
- 407 T. Luong, H. Pham, and C. D. Manning. Effective approaches to attention-based neural machine
408 translation. In *EMNLP*, 2015.
- 409 S. Merity, C. Xiong, J. Bradbury, and R. Socher. Pointer sentinel mixture models. *ICLR*, 2017.
- 410 R. Nallapati, B. Zhou, C. dos Santos, G. glar Gulcehre, and B. Xiang. Abstractive text summarization
411 using sequence-to-sequence rnn and beyond. *CoNLL2016*, page 280, 2016.
- 412 A. Neelakantan, Q. V. Le, M. Abadi, A. McCallum, and D. Amodei. Learning a natural language
413 interface with neural programmer. In *ICLR*, 2017.
- 414 E. W. Ngai, L. Xie, and D. C. Chan. Application of data mining techniques in customer relationship
415 management: A literature review and classification. *Expert systems with applications*, 36(2):
416 2592-2602, 2009.
- 417 P. Pasupat and P. Liang. Compositional semantic parsing on semi-structured tables. In *ACL*, 2015.
- 418 R. Puri, C. Xiong, and R. Socher. A deep reinforced model for abstractive summarization. *arXiv*
419 *preprint arXiv:1703.04304*, 2017.
- 420 J. Pennington, R. Socher, and C. D. Manning. Glove: Global vectors for word representation. In
421 *EMNLP*, 2014.
- 422 A.-M. Papescu, O. Etzioni, and H. Kautz. Towards a theory of natural language interfaces to databases.
423 In *Proceedings of the 8th International Conference on Intelligent User Interfaces*, pages 149-157,
424 ACM, 2003.
- 425 P. J. Price. Evaluation of spoken language systems: The ATIS domain, 1990.
- 426 S. Reddy, M. Lapata, and M. Steedman. Large-scale semantic parsing without question-answer pairs.
427 *TACL*, 2:377-392, 2014.

10

20

30

WO 2018/213530

33

PCT/US2018/033099

- 432 J. Schultson, N. Heess, T. Weber, and P. Abbeel. Gradient estimation using stochastic computation
433 graphs. In *NIPS*, 2015.
- 434 M. J. Seo, A. Kembhavi, A. Farhadi, and H. Hajishirzi. Bidirectional attention flow for machine
435 comprehension. *ICLR*, 2017.
- 436 N. Srivastava, G. E. Hinton, A. Krizhevsky, I. Sutskever, and R. Salakhutdinov. Dropout: a simple
437 way to prevent neural networks from overfitting. *Journal of Machine Learning Research*, 15:
438 1935–1958, 2014.
- 439 J. Slone and D. Mladenic. Joint learning of ontology and semantic parser from text. *Intelligent Data
440 Analysis*, 21:19–38, 2017.
- 441 I. Sutskever, O. Vinyals, and Q. V. Le. Sequence to sequence learning with neural networks. In *NIPS*,
442 2014.
- 443 L. R. Tang and R. J. Mooney. Using multiple clause constructors in inductive logic programming for
444 semantic parsing. In *ECML*, 2001.
- 445 O. Vinyals, M. Pottoruto, and N. Jaitly. Pointer networks. In *NIPS*, 2015.
- 446 Y. Wang, J. Bernat, and P. Liang. Building a semantic parser overnight. In *ACL*, 2015.
- 447 Y. W. Wong and R. J. Mooney. Learning synchronous grammars for semantic parsing with lambda
448 calculus. In *ACL*, 2007.
- 449 C. Xiong, V. Zhong, and R. Socher. Dynamic coercion networks for question answering. *ICRL*,
450 2017.
- 451 J. M. Zelle and R. J. Mooney. Learning to parse database queries using inductive logic programming.
452 In *AAAI/AAJ*, vol. 2, 1996.
- 453 L. S. Zettlemoyer and M. Collins. Learning to map sentences to logical form: Structured classification
454 with probabilistic categorical grammars. In *Uncertainty in Artificial Intelligence*, 2003.
- 455 L. S. Zettlemoyer and M. Collins. Online learning of reduced csg grammars for parsing to logical
456 form. In *EMNLP-CoNLL*, 2007.

10

20

30

WO 2018/213530

34

PCT/US2018/033099

A Collection of WikisQL

WikisQL is collected in a paraphrase phases as well as a verification phase. In the paraphrase phase, we use tables extracted from Wikipedia by Bhagavata et al [Bhagavata et al., 2013] and remove small tables according to the following criteria:

- the number of cells in each row is not the same
- the content in a cell exceed 50 characters
- a header cell is empty
- the table has less than 5 rows or 5 columns
- over 10% of the cells of a row contain identical content

We also remove the last row of a table because a large quantity of HTML tables tend to have summary statistics in the last row and hence the last row does not adhere to the table schema defined by the header row.

For each of the table that passes the above criteria, we randomly generate 5 SQL queries according to the following rules:

- the query follows the format `SELECT agg_op agg_col from table where cond1_col cond1_op cond1 AND cond2_col cond2_op cond2 ...`
- the aggregation operator `agg_op` can be empty or `COUNT`. In the event that the aggregation column `agg_col` is numeric, `agg_op` can additionally be one of `MAX` and `MIN`
- the condition operator `cond_op` is `=`. In the event that the corresponding condition column `cond_col` is numeric, `cond_op` can additionally be one of `>` and `<`
- the condition `cond` can be any possible value present in the table under the corresponding `cond_col`. In the event that `cond_col` is numerical, `cond` can be any numerical value sampled from the range from the minimum value in the column to the maximum value in the column.

We only generate queries that produce a non-empty result set. To enforce succinct queries, we remove conditions from the generated queries if doing so does not change the execution result.

For each query, we generate a crude question using a template and obtain a human paraphrase via crowdsourcing on Amazon Mechanical Turk. In each Amazon Mechanical Turk HIT, a worker is shown a table as well as its generated questions and asked to provide paraphrases for each question.

After obtaining natural language interferences from the paraphrase phase, we give each question-paraphrase pair to two other workers in the verification phase to verify that the paraphrase and the original question contain the same meaning.

We then filter the initial collection of paraphrases using the following criteria:

- the paraphrase must be deemed correct by at least one worker during the verification phase
- the paraphrase must be sufficiently different from the generated question, with a character-level edit distance greater than 10

Figure 2 shows one example from WikisQL. Here, the corresponding generated question is "what is the total number of CPL Team where College is York".

565 B: Attentional Seq2Seq Baseline

566 For our baseline, we employ the attentional sequence-to-sequence model, which has led to significant
 567 improvements in machine translation [Bahdanau et al., 2015] and neural semantic parsing [Dong
 568 and Lapata, 2016]. The model is based on the global attention encoder-decoder model (with input
 569 feeding) described by Luong et al. [2015]. We use an implementation based on OpenSMT [Klein
 570 et al.].

571 We use the same two-layer, bidirectional, stacked LSTM encoder as described previously. The
 572 decoder is a two-layer, unidirectional, stacked LSTM. The decoder LSTM is almost identical to that
 573 described by Equation 2, with the sole difference coming from input feeding.

$$y_t = \text{LSTM}([\text{emb}(y_{t-1}); h_{t-1}^{\text{dec}}], y_{t-1}) \quad (8)$$

574 where the d_{att} -dimensional h_t^{att} is the attentional context over the input sequence during the t th
 575 decoding step, computed as:

$$o_{t,j}^{\text{att}} = h_t^{\text{dec}} (W^{\text{att}} h_j^{\text{enc}})^T \quad \beta_t^{\text{att}} = \text{softmax}(o_t^{\text{att}}) \quad (9)$$

$$h_t^{\text{att}} = \sum_j \beta_{t,j} h_j^{\text{enc}} \quad (10)$$

576 To produce the output token during the t th decoder step, the concatenation of the decoder state and
 577 the attention context is given to a final linear layer to produce a distribution o_t^{out} over words in the
 578 target vocabulary:

$$o_t^{\text{out}} = \text{softmax}(U^{\text{out}}(h_t^{\text{dec}} \parallel h_t^{\text{att}})) \quad (11)$$

579 where U^{out} has dimensions $N^{\text{vocab}} \times d_{\text{att}}$ and N^{vocab} is the size of the output vocabulary.

580 During training, teacher forcing is used. During inference, a beam size of 5 is used and generated
 581 unknown words are replaced by the input words with the highest attention weight.

WO 2018/213530

36

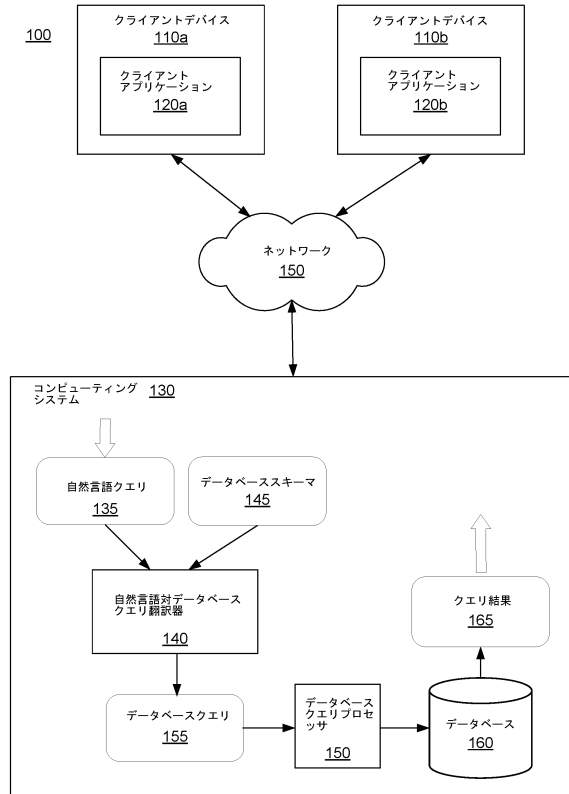
PCT/US2018/033099

Fig. C Predictions by Seq2SQL

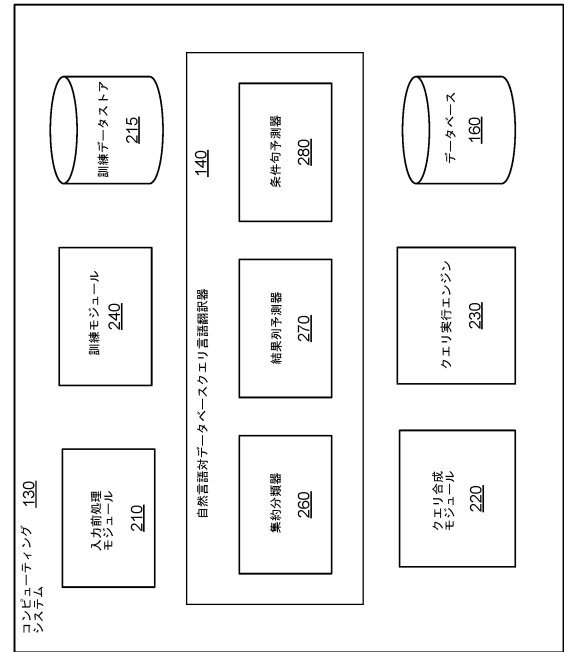
Q	when connecticut & villanova are the regular season winner how many consecutive years (in) are they?
P	SELECT COUNT(*) tournament, player, city, WHERE regular season winner city = connecticut & villanova
S	SELECT COUNT(*) tournament, winner, city, WHERE tournament, winner = connecticut & villanova
G	SELECT COUNT(*) tournament, winner, city, WHERE regular season winner = connecticut & villanova
Q	what are the aggregate scores of those teams when the GOLF results are 0-1?
P	SELECT aggregate WHERE loc = 0-1
S	SELECT COUNT(*) agg, score, WHERE loc lag = 0-1
G	SELECT agg, score, WHERE loc lag = 0-1
Q	what is the race name of the 12th round robin, how many are there by, first and the race position?
P	SELECT race name WHERE location = 12th AND race position = 0-1, first, race jersey AND
S	SELECT race name WHERE race = 12 AND race = 0-1, first AND race position = 0-1, first
G	SELECT race name WHERE race = 12 AND race position = 0-1, first
Q	what city is in 1900?
P	SELECT city WHERE frequency = 1900
S	SELECT city of license WHERE frequency = 1900
G	SELECT city of license WHERE frequency = 1900
Q	how many were there by long used in the social party?
P	SELECT AND party = socialist
S	SELECT COUNT(*) the name where the name = socialist
G	SELECT the name WHERE the name = socialist
Q	what year did a Plymouth vehicle leave February 9?
P	SELECT year that WHERE date = February 9 AND race time = Plymouth 9
S	SELECT year that WHERE date = Plymouth 9 AND race time = February 9
G	SELECT year that WHERE manufacturer = Plymouth AND date = February 9

Table 4: Examples predictions by the models on the dev split. Q denotes the natural language question and G denotes the corresponding ground truth query. P, S, and G denote, respectively, the queries produced by the Augmented Pointer Network, Seq2SQL, without reinforcement learning, Seq2SQL. We omit the FROM table part of the query for succinctness.

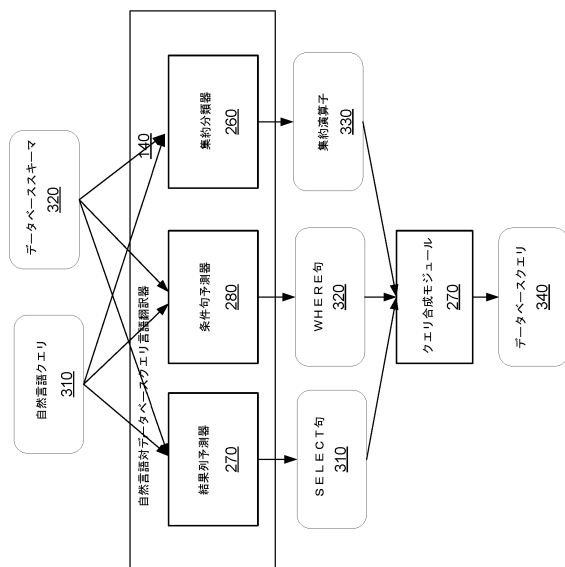
【図 1】



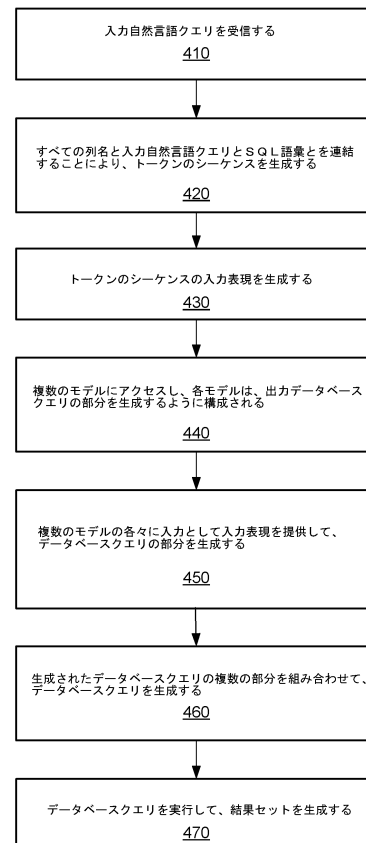
【図 2】



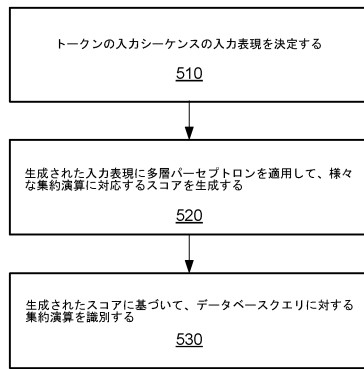
【図 3】



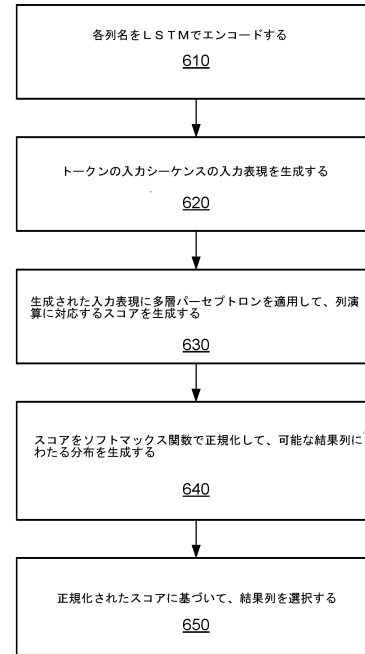
【図 4】



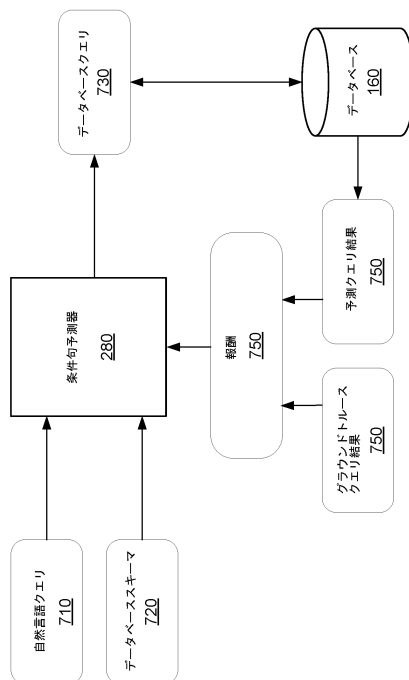
【図 5】



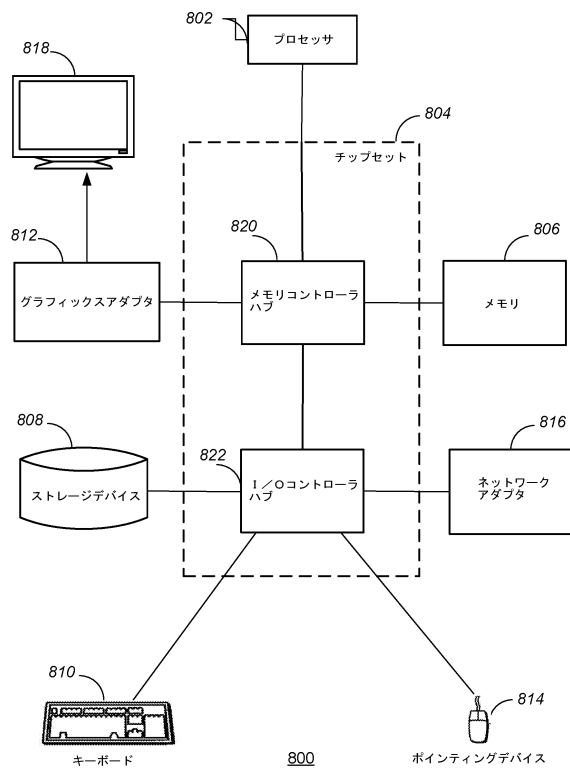
【図 6】



【図 7】



【図 8】



フロントページの続き

(72)発明者 ゾン, ヴィクター

アメリカ合衆国 カリフォルニア州 94105, サンフランシスコ, ザ ランドマーク アット
ワン マーケット ストリート, スイート 300, セールスフォース ドット コム インコ
ーポレイティッド

(72)発明者 ション, カイミング

アメリカ合衆国 カリフォルニア州 94105, サンフランシスコ, ザ ランドマーク アット
ワン マーケット ストリート, スイート 300, セールスフォース ドット コム インコ
ーポレイティッド

(72)発明者 ソーチャー, リチャード

アメリカ合衆国 カリフォルニア州 94105, サンフランシスコ, ザ ランドマーク アット
ワン マーケット ストリート, スイート 300, セールスフォース ドット コム インコ
ーポレイティッド

審査官 多胡 滋

(56)参考文献 国際公開第2016/151690(WO, A1)

米国特許出願公開第2013/0239006(US, A1)

(58)調査した分野(Int.Cl., DB名)

G06F 16/24

G06N 3/00