



Offre n°2024-07921

Doctorant F/H Des Grands Modèles de Language pour la détection et la correction des erreurs dans les applications HPC

Type de contrat : CDD

Niveau de diplôme exigé : Bac + 5 ou équivalent

Fonction : Doctorant

A propos du centre ou de la direction fonctionnelle

Le centre Inria de l'université de Bordeaux est un des neuf centres d'Inria en France et compte une vingtaine d'équipes de recherche. Le centre Inria est un acteur majeur et reconnu dans le domaine des sciences numériques. Il est au cœur d'un riche écosystème de R&D et d'innovation : PME fortement innovantes, grands groupes industriels, pôles de compétitivité, acteurs de la recherche et de l'enseignement supérieur, laboratoires d'excellence, institut de recherche technologique...

Contexte et atouts du poste

Nous proposons un contrat de thèse sur une durée de 3 ans dans l'équipe Storm (<https://team.inria.fr/storm/>) du centre Inria de l'Université de Bordeaux.

Mission confiée

Afin de résoudre les plus grands problèmes scientifiques en un temps raisonnable, les applications sont parallélisées et lancées sur des supercalculateurs. Cependant, ces supercalculateurs sont de plus en plus complexes et puissants, ce qui entraîne une évolution des applications (ex., nouveaux algorithmes pour le passage à l'échelle, combinaison de modèles de programmation parallèle). Cette évolution implique de nombreux défis de programmation et un réel besoin d'outils et techniques pour aider les développeurs à utiliser au mieux les différentes machines et architectures à leur disposition. En effet, à grande échelle, les développeurs d'applications font face à de nouvelles erreurs, liées au parallélisme, souvent difficiles à analyser et corriger. Aujourd'hui, s'assurer que les applications parallèles s'exécutent correctement devient aussi important que d'obtenir de bonnes performances.

Les grands modèles de langage (LLMs) sont un sujet de recherche en pleine évolution. En particulier, leurs récents succès pour générer du texte pertinent et répondre à des questions en font des candidats attrayants dans le domaine de la vérification.

Objectif:

L'objectif de cette thèse est d'exploiter et d'adapter les Grands Modèles de Language pour identifier et corriger les erreurs dans les programmes parallèles. Pour cela, nous proposons d'entraîner des modèles sur des ensembles de données soigneusement générés et étiquetés grâce à une combinaison de techniques d'apprentissage et de traitement du langage naturel.

Collaboration :

La personne recrutée sera sous la direction d'Emmanuelle Saillard et Mihail Popov. Elle sera également en lien avec Pablo Oliveira (Université de Versailles) et Eric Petit (Intel).

Principales activités

Le programme de recherche est découpé en 4 axes d'exploration.

Axe 1 : Cr  ation d'un jeu de donn  es

Un jeu de donn  es de haute qualit   est une condition n  cessaire pour cr  er des mod  les pr  cis. Pour cr  er notre jeu de donn  es, nous nous appuierons sur deux sources compl  mentaires contenant des codes corrects et incorrects. Dans un premier temps, nous exploiterons la base de donn  es git d'EasyPAP [1], une plateforme qui enseigne la programmation parall  le. Bien que limit  e en taille, le code soumis par les   tudiants est repr  sentatif des erreurs que font les d  butants. Nous explorerons ensuite Github via son API int  gre  e pour collecter des codes r  els et plus conséquents en taille. Les projets seront s  lectionn  s selon les issues, pull requests et descriptions des commits. Nous r  cup  rerons le code avant et apr  s les commits pertinents.

Axe 2 : Labellisation

Une fois le jeu de donn  es cr    , l'  tape cruciale est d'  tiqueter les programmes, c'est-  dire d'associer chaque programme avec un label (erreur pr  sente dans le code ou corrig  e). Pour cela, on utilisera des techniques de NLP. Les descriptions des commits et toute m  ta-information associ  e (e.g., CI) seront analys  es avec TF-IDF (ou optionnellement des textes d'embeddings    la word2vec). Les vecteurs obtenus seront trait  s avec NMF [2] pour en extraire les diff  rentes classes d'erreurs que nous   tudierons. En parall  le, nous pourrions   galement directement utiliser des LLMs (e.g., ChatGPT) sur les commit pour les grouper. De plus, nous analyserons l'embedding des codes avant et apr  s chaque commit [3] : les vecteurs obtenus seront clusteris  s pour grouper des changements similaires.    terme, nous unifions les deux classifications pour cr  er un processus de labellisation plus g  n  ral.

Axe 3 : Entra  nement des mod  les

Nous visons deux types de mod  les. Nous commencerons par cr  er des mod  les supervise  s (Code2Error) qui prennent le code source (ou une repr  sentation du compilateur, e.g., LLVM IR) d'un programme et pr  disent la cat  gorie d'erreur associ  e au programme (base  e sur la labellisation). Ces mod  les permettront de classer les codes incorrects et d'enrichir les descriptions des probl  mes. En d  tail, Code2Error utilisera un embedding (e.g., ir2vec, code2vec) pour g  n  rer des vecteurs,    partir des codes, auxquels nous appliquerons un mod  le supervise   (e.g., arbre de d  cision) pour d  cider du label. Les codes avant et apr  s le commit serviront    donner    l'arbre la version incorrecte et sa correction. Nous avons d  j   valid   une version pr  liminaire (ir2vec & arbre de decision) sur 2000 codes tests d  di  s pour la v  rification MPI et souhaitons passer ce mod  le    l'  chelle sur de vrais codes.

Ensuite, nous utiliserons des LLMs (Code2Fix). Pour chaque erreur (et donc groupe de commits associ  s), nous entra  nerons un LLM sp  cialis  . Ce LLM recevra les codes corrects et incorrects associ  s    une certaine erreur. Nous utiliserons ici les codes sources (car plus utile pour l'utilisateur) et entra  nerons (fine tuning) le LLM pour passer de la version erron  e    la version correcte. Nous pourrions appliquer Code2Error sur un programme inconnu pour identifier le type d'erreur qu'il contient, et appeler le LLM Code2Fix associ      l'erreur pour essayer de la r  soudre. Notre intuition est qu'un LLM sp  cialis   par erreur sera plus efficace. Enfin, on pourra explorer la granularit   du code pour Code2Error & Code2Fix. De petites granularit  s seront faciles    g  rer pour le mod  le et donc pour trouver la localisation de l'erreur au moment de la correction mais pourront manquer de contexte pour traiter certaines erreurs compliqu  es. Ce sera un compromis    explorer.

Axe 4 : Disse  mination

Les diff  rents mod  les (Code2Fix, Code2Error) seront appliqu  s    des projets existants pour chercher et corriger des erreurs existantes. Nous validerons   galement nos mod  les sur des erreurs que nous aurons exclues du jeu de donn  es pour l'apprentissage afin de mettre en avant la g  n  ralisation de notre m  thode et estimer    quel point deux erreurs sont similaires (si nous pouvons pr  dire une erreur avec des informations provenant d'une autre erreur, il est probable qu'elles soient li  es). Les experts en outils de v  rification pourront utiliser cette information pour d  finir de nouvelles topologies d'erreurs. Enfin, nous envisageons d'  tendre notre ensemble de donn  es avec de nouveaux codes g  n  r  s automatiquement par le biais des LLMs (Dataset2Code).

R  f  rences :

[1] A. Lasserre, R. Namyst, and P.-A. Wacrenier. Easypap : a framework for learning parallel programming. In 2020 IEEE International Parallel and Distributed Processing Symposium Workshops (IPDPSW), pages 276– 283, 2020.

[2] S. Heldens, P. Hijma, B. Werkhoven, J. Maassen, A. Belloum, and R. Van Nieuwpoort. The landscape of exascale research : A data-driven literature analysis. ACM Computing Surveys, 53(2) :1–43, Mar. 2020.

[3] H. Wang, G. Ye, Z. Tang, S. H. Tan, S. Huang, D. Fang, Y. Feng, L. Bian, and Z. Wang. Combining graph-based learning with automated data collection for code vulnerability detection. Trans. Info. For. Sec., 16 :1943–1958, jan 2021.

Comp  tences

- Motivation
- Curiosit   et capacit      apprendre de nouveaux concepts

- Expérience avec l'écriture de scripts (ex., Python)
- Maîtrise des bases Linux
- Des connaissances en ML est un plus

Avantages

- Restauration subventionnée
- Transports publics remboursés partiellement
- Congés: 7 semaines de congés annuels + 10 jours de RTT (base temps plein) + possibilité d'autorisations d'absence exceptionnelle (ex : enfants malades, déménagement)
- Possibilité de télétravail et aménagement du temps de travail
- Équipements professionnels à disposition (visioconférence, prêts de matériels informatiques, etc.)
- Prestations sociales, culturelles et sportives (Association de gestion des œuvres sociales d'Inria)
- Accès à la formation professionnelle
- Sécurité sociale

Rémunération

Montant salaire brut 1e année : 2100€

Montant salaire brut 2e et 3e année : 2190€

Informations générales

- **Thème/Domaine** : Calcul distribué et à haute performance
Calcul Scientifique (BAP E)
- **Ville** : Talence
- **Centre Inria** : [Centre Inria de l'université de Bordeaux](#)
- **Date de prise de fonction souhaitée** : 2024-10-01
- **Durée de contrat** : 3 ans
- **Date limite pour postuler** : 2024-07-31

Contacts

- **Équipe Inria** : [STORM](#)
- **Directeur de thèse** :
Saillard Emmanuelle / emmanuelle.saillard@inria.fr

A propos d'Inria

Inria est l'institut national de recherche dédié aux sciences et technologies du numérique. Il emploie 2600 personnes. Ses 215 équipes-projets agiles, en général communes avec des partenaires académiques, impliquent plus de 3900 scientifiques pour relever les défis du numérique, souvent à l'interface d'autres disciplines. L'institut fait appel à de nombreux talents dans plus d'une quarantaine de métiers différents. 900 personnels d'appui à la recherche et à l'innovation contribuent à faire émerger et grandir des projets scientifiques ou entrepreneuriaux qui impactent le monde. Inria travaille avec de nombreuses entreprises et a accompagné la création de plus de 200 start-up. L'institut s'efforce ainsi de répondre aux enjeux de la transformation numérique de la science, de la société et de l'économie.

L'essentiel pour réussir

Le ou la candidate doit avoir un bon niveau de programmation.

De plus, la personne recrutée devra relire de la bibliographie scientifique, écrire des rapports/articles et présenter ses travaux devant la communauté. De ce fait, un bon niveau de communication en anglais sera fortement apprécié.

Attention: Les candidatures doivent être déposées en ligne sur le site Inria. Le traitement des candidatures adressées par d'autres canaux n'est pas garanti.

Consignes pour postuler

Si vous êtes intéressés, merci de bien vouloir candidater via le site [jobs.inria](https://jobs.inria.fr) avec les documents suivants :

- CV
- lettre de motivation
- notes master
- lettre de recommandation le cas échéant

Sécurité défense :

Ce poste est susceptible d'être affecté dans une zone à régime restrictif (ZRR), telle que définie dans le décret n°2011-1425 relatif à la protection du potentiel scientifique et technique de la nation (PPST). L'autorisation d'accès à une zone est délivrée par le chef d'établissement, après avis ministériel favorable, tel que défini dans l'arrêté du 03 juillet 2012, relatif à la PPST. Un avis ministériel défavorable pour un poste affecté dans une ZRR aurait pour conséquence l'annulation du recrutement.

Politique de recrutement :

Dans le cadre de sa politique diversité, tous les postes Inria sont accessibles aux personnes en situation de handicap.