

오픈소스 LLM을 위한 문서 유형별 파인튜닝 적용 방법론 : Unsloth CPT, SFT

이용구*, 김민규**

한국수력원자력 중앙연구원

e-mail: lee.yongku@khnp.co.kr, kim.minkyu@khnp.co.kr

Document Type-Specific Fine-Tuning Methodology for Open Source LLM : Unsloth CPT, SFT

Yong-Ku Lee*, Min-Kyu Kim**

*,**Central Research Institute, Korea Hydro & Nuclear Power Co., Ltd.

요 약

본 연구에서는 기업 내부 데이터의 문서 유형에 따라 오픈소스 대형언어모델(LLM)을 효과적으로 파인튜닝하는 이중 접근법을 제안한다. 구체적으로, 정형적이고 일방향적인 문서(예: 법규, 지침, 사내 규정 등)에 대해서는 연속적인 추가 학습(Continued Pre-Training, CPT)을 수행하여 모델의 도메인 지식을 강화하고, 양방향 이해가 필요한 문서(예, 업무 계획서, 품의서 등)에 대해서는 Supervised Fine-Tuning(SFT)을 적용하여 해당 맥락에서의 모델 응답 능력을 학습시키는 데에 목적이 있다. 파인튜닝 대상 모델로는 Qwen2.5와 Gemma2 등의 공개 LLM을 활용하였으며, 경량 파인튜닝에 최적화된 Unsloth 프레임워크를 통해 효율적인 학습 절차를 구현하였다. 본 논문에서는 CPT와 SFT 방법론의 차이점과 문서 유형별 적요기준을 정의하고, Unsloth기반의 모델 학습 절차 및 데이터 전처리 전략을 논의한다. 이를 통해 데이터 구성의 방식과 실무 적용 관점에서 이러한 방법론의 유용성을 탐색하였으며, 제안된 접근법이 향후 기업 내부 데이터를 활용한 LLM 성능 향상에 기여할 것으로 기대된다.

1. 서론

대형언어모델(LLM)은 GPT-4, Bard 등과 같이 범용적인 지식에 기반한 강력한 언어 처리 능력으로 다양한 분야에 활용되고 있다. 그러나 이러한 범용 LLM을 기업 내부 업무에 직접 활용하는 데에는 몇 가지 한계가 존재한다. 우선, 최신 거대 모델들은 대부분 API 형태로 제공되는 폐쇄형 모델인 경우가 많아, 회사의 민감한 내부 데이터를 외부 서비스로 보내야 하는 보안 및 프라이버시 이슈가 있다. 또한 거대 모델의 크기와 복잡성으로 인해 온프레미스 환경에서의 사용이 어려워, 인터넷 연결 없이 로컬에서 실행하거나 신속하게 커스터마이징하기 어렵다. 이러한 이유로 기업들은 자체 오픈소스 LLM을 도입하여 내부 데이터로 파인튜닝하고, 모델을 사내에 배포하는 방향을 모색하고 있다. 이는 외부 의존성을 낮추고, 내부 도메인 지식을 모델에 통합할 수 있다는 장점이 있다. 기업 내부에는 법령, 지침, 내부 규정, 업무 보고서, 계획서 등 다양한 형태의 문서 데이터가 축적되어 있으며, 이를 LLM에 학습시킬 경우 사내 검색, 업무 자동화, 의사결정 지원 등에서 지능형 활용이 가능해진다. 하지만 범용 사전학습된 LLM

은 이러한 특수 도메인의 세부 지식이나 맥락을 충분히 반영하지 못할 수 있으므로, 추가 파인튜닝을 통해 모델을 기업 환경에 맞게 도메인 적응시키는 과정이 필수적이다. 특히 내부 문서는 그 유형과 역할에 따라 내용 구조와 활용 방식이 상이하여, 문서별 특성을 고려한 맞춤형 파인튜닝 전략이 요구된다. 예를 들어, 법규나 사내 규정과 같은 문서는 주로 일방향으로 서술된 정형적 정보의 집합이며, 모델 입장에서는 이러한 문서를 통하여 사실적 지식과 전문 용어를 흡수하는 것이 중요하다. 이처럼 지식 축적이 목적인 경우에는 대량의 텍스트를 통해 언어모델을 추가로 사전훈련하는 접근이 유효하다. 반면에, 업무 계획서나 품의서와 같은 문서는 작성자와 결재자 간의 상호작용 맥락을 포함하고 있으며, 문서에 담긴 의도 파악, 요약 및 질의 응답 등의 이해 및 생성 능력이 요구된다. 이러한 문서에 대해 모델을 활용하려면 단순히 내용을 기억하는 것을 넘어, 주어진 문서를 바탕으로 적절한 응답이나 산출물을 생성하는 능력을 길러야 한다. 따라서 문서의 유형에 따라 파인튜닝 목표를 달리 설정하고 그에 상응하는 기법을 적용하는 것이 합리적이다. 본 연구에서는 기업 문서를 활용한 LLM 파인튜닝에 있어 문서 유형별 이중 접근법을 제안한다. 구체적으로, 앞서 언급한 바와 같이 정형·일방향 문서에는 CPT(Continued Pre-Training) 방식의 비지도 사전학습을 통해 도메인 지식을 주입하고, 양방향 상호작용 문서에는

****SFT(Supervised Fine-Tuning)****를 통해 문서 이해 및 응답 생성에 대한 지도 학습을 수행한다. 이를 위해 파인튜닝 대상 모델로 Qwen2.5와 Gemma2 등을 선정하였다. Qwen2.5는 Alibaba Cloud에서 공개한 최신 LLM 시리즈로 수십억~수백억 매개변수 규모의 모델을 포함하며, 다국어 및 수리·코딩 능력에서 우수한 성능을 보이는 것으로 보고되고 있다. Gemma2는 Google에서 공개한 경량화된 LLM 패밀리로 2B, 9B, 27B 등 다양한 크기로 제공되어 실용성을 높인 모델이다. 이러한 최신 공개 모델들을 기반으로 하면 기업 데이터에 특화하면서도 기본 언어 능력이 우수한 기반을 확보할 수 있다. 더욱이, 오픈소스 모델을 선택함으로써 모델 파인튜닝 및 서비스 과정을 모두 사내 환경에서 실행할 수 있어 데이터 유출 위험을 낮추고 제어권을 확보할 수 있다. 파인튜닝 방법론을 구현하는 데에는 Unsloth라는 경량화 학습 프레임워크를 활용하였다. Unsloth는 LoRA 및 QLoRA와 같은 파라미터 효율 기법을 사용하여 대형 모델의 파인튜닝 시 필요한 메모리와 연산량을 크게 감소시킨다. 이를 통해 일반적인 풀 파라미터 튜닝 대비 25배 빠른 학습 속도와 약 70% 적은 메모리 사용으로 LLM을 미세조정할 수 있다고 보고되어 있다. 본 연구에서도 Unsloth를 사용하여 제한된 GPU 자원(예: 단일 Tesla T4 등)으로 Qwen2.5(수억수십억 규모)와 Gemma2 모델에 대한 CPT와 SFT 단계를 수행하였다. 이처럼 경량 파인튜닝 환경을 활용하면, 기업 내 보편적인 하드웨어로도 자체 LLM 튜닝이 가능해져 실무 적용성을 높일 수 있다. 본 연구는 방법론적 탐색에 초점을 맞추고 있으며, 제안된 접근법의 유효성 개념 증명에 중점을 둔다. 따라서 모델 간 성능 정량 비교 실험은 생략하고, 대신 데이터 준비 과정의 난이도, 파인튜닝 절차의 실용성, 그리고 기업 환경에서 예상되는 이점과 한계를 논의한다. 이러한 접근을 통해, 제한된 데이터와 자원으로 도메인 특화 LLM을 구축하기 위한 현실적인 방안을 모색하고자 한다.

2. 본론

2.1 CPT와 SFT의 개요 및 적용 기준

2.1.1 Continued Pre-Training(CPT)는 사전학습된 언어모델에 도메인 특화 코퍼스를 추가로 학습시켜 지식을 연속적으로 확장하는 방법이다. 구체적으로, 모델이 다음 토큰을 예측하도록 하는 언어모델링 목표를 유지한 채로 새로운 텍스트 데이터를 계속 학습시킴으로써, 해당 텍스트에 담긴 용어와 사실을 모델 파라미터에 내재화한다. 본 연구에서는 기업 내부의 일방향 문서들을 CPT를 통해 모델에 주입하였다. 예를 들면, 사내 규정, 법률 조항, 업

무 지침서, 기술 매뉴얼과 같은 문서들은 작성자가 정보를 전달하기 위해 쓰인 것으로, 모델이 이를 그냥 읽어내려가며 학습하는 것만으로도 유용한 지식을 얻을 수 있다. 이러한 문서에 CPT를 적용하면 모델이 기업 고유의 용어 체계, 도메인 배경지식, 문서 특유의 문체 등을 습득하게 된다. CPT의 장점은 별도의 레이블된 출력이 필요 없으므로 이미 존재하는 방대한 텍스트를 그대로 활용할 수 있다는 점이다. 즉, 사람이 직접 질문-답 형식의 학습 데이터를 만들지 않아도 되므로 데이터 준비 부담이 적고, 확보된 모든 문서를 최대한 활용할 수 있다. 다만 CPT는 모델에게 새로운 정보를 그냥 기억하도록 하는 것이므로, 이를 통해 모델이 문제를 해결하는 방법을 학습하는 것은 아니다. 지식 축적이 주 목적인 CPT는, 문서 자체를 이해하고 표현할 수 있게 하지만 그 지식을 어떻게 활용할지는 별도의 과정이 필요할 수 있다.

2.1.2 Supervised Fine-Tuning (SFT)은 모델에 특정 작업 수행 능력을 가르치기 위해 정답이 주어진 예제들로 추가 학습을 진행하는 기법이다. 일반적으로 사전학습이나 CPT로 기본기를 쌓은 모델에 대해, 사람이 설계한 입력-출력 쌍(지도 데이터, QA 데이터셋)을 사용하여 모델이 주어진 입력에 적절한 응답을 생성하도록 미세조정한다. 본 연구에서는 양방향 이해가 필요한 문서들에 SFT를 적용하였다. 예컨대, 업무 계획서나 품의서는 단순 정보 나열이 아니라 맥락과 의도가 담긴 서술이며, 이를 활용한 작업(예: 요약, 질의응답, 검토 코멘트 생성 등)은 단순 암기 이상의 능력을 요구한다. 모델이 이러한 작업을 수행하려면 문서의 내용을 완전히 이해하고, 추가적인 추론이나 요약을 해야 한다. SFT 단계에서는 이러한 능력을 길러주기 위해, 해당 문서를 기반으로 한 질문-응답 쌍, 요약문 생성 예시, 지침에 따른 응답 생성 예시 등의 훈련 데이터를 구성한다. 예를 들어 한 프로젝트 계획서에 대해 "이 계획의 주요 목표는 무엇인가?"라는 질문을 만들고, 그에 대한 정답을 계획서 내용에서 발췌하거나 요약하여 제공함으로써 모델을 훈련시킬 수 있다. 또는 품의서의 본문을 주고 "해당 품의 내용을 한 문단으로 요약하시오"라는 지시와 모범 답안을 함께 제시하여 모델이 요약 작업을 배우도록 할 수 있다. 이러한 SFT 데이터셋을 구축함으로써, 모델은 단순히 문장을 예측하는 것을 넘어 문제를 해결하거나 응답을 구성하는 패턴을 익히게 된다. SFT의 단점은 필요한 지도 데이터의 생성 비용이 높다는 점인데, 기업 문서에 특화된 Q&A나 요약 예시를 만들기 위해서는 도메인 전문가의 노력이나 추가적인 데이터 생성 알고리즘이 필요하다. 따라서 어떤 문서에 SFT를 적용할지는 비용 대비 효과를 고려하여 결정해야 한다. 정형 지식 문서에 대해서는 굳이 질문-답변 형태로 변환하지 않아도 CPT로 충분한 반면, 상호작용형 문서는 SFT를 통해서만 모델이 실제 활용 가능한 출력을

내도록 훈련될 수 있기 때문이다.

이상의 내용을 정리하면, 문서의 성격에 따라 적용되는 파인튜닝 전략은 다음과 같다.

2.1.4 CPT 대상은 일방향 지식 전달 문서로써 기업 내부 규정, 법률/법규 조문, 기술 지침서, 제품 매뉴얼 등. 이러한 문서는 주로 사실과 지식을 전달하므로 모델에 그대로 읽혀서 축적하게 함으로써 효과를 얻는다. (예: 모델이 회사 인사규정을 통해로 학습하면, 추후에 "육아휴직 기간은 최대 몇 년인가?"와 같은 질문에 근거를 가지고 답할 수 있게 됨)

2.1.5 SFT 대상은 양방향 이해 및 생성 문서로써 사업 제안서, 프로젝트 계획서, 품의서, 회의록 등. 이러한 문서는 내용을 토대로 추가적인 행동(요약, 질의응답, 의견 생성)이 이루어지므로, 해당 작업에 대한 지도 학습으로 모델의 응답 생성 능력을 길러 준다. (예: 모델이 여러 건의 사업 제안서와 그에 대한 요약본으로 SFT 되었다면, 새로운 제안서를 입력받았을 때 자동으로 핵심 요점을 요약해줄 수 있음)

이처럼 CPT는 지식 주입(Knowledge Injection), SFT는 특정 과제 학습으로 구분되며, 두 접근법을 상보적으로 활용하는 것이 기업 문서 분야에 특화된 LLM을 만드는 핵심 전략이다.

2.2 Unsloth 기반 파인튜닝 구현

2.2.1 Unsloth Training Framework 특징

제안하는 CPT와 SFT 방법론을 실제 모델에 적용하기 위해, 본 연구에서는 Unsloth 프레임워크를 사용하였다. Unsloth는 공개 LLM들의 파인튜닝 작업을 간소화하고 효율화하기 위해 개발된 도구로, 내부적으로 LoRA (Low-Rank Adaptation) 및 QLoRA (4-bit 양자화 LoRA) 등의 기술을 활용한다. 이를 통해 대용량 모델도 모든 파라미터를 업데이트하지 않고 소량의 추가 파라미터만 학습하거나, 모델을 저장밀도로 메모리 효율적으로 로드하여 학습할 수 있다. 예를 들어, Unsloth를 사용하면 LLaMA, Mistral, Gemma, Qwen 등의 수십억~수천억 규모 LLM을 단일 GPU 환경에서 튜닝할 수 있으며, 실제로 8비트 또는 4비트 양자화를 통해 메모리 사용량을 절감하고 학습 속도를 높일 수 있다고 보고되어 있다. 특히 Gemma-2 9B 모델의 경우 Unsloth를 활용하면 기존 대비 2.4배 빠른 학습과 58%의 메모리 절감 효과를 보이는 등, 자원 제약 하에서도 실용적인 파인튜닝이 가능함이 알려져 있다.

2.2.2 CPT 학습절차 구성

먼저 CPT 단계에서는 Unsloth의 UnslothTrainer를 사용하며, Continued Pre-Trainig을 안정적으로 진행하기 위하여 모델 로드 시 lm_head, embed_tokens 모듈을 포함해 학습되도록 설정해야 한다. 이를 통해 기존 사전학습된 가중치 뿐만 아니라 언어모델 출력부와 임베딩 레이어까지 함께 학습하는 것이 가능해진다. CPT 학습을 위한 말뭉치는 선정된 정형 문서들의 텍스트로 이루어져 있으며, 이를 Unsloth에 입력하여 모델이 연속된 문맥에서 해당 문서들을 읽으며 확률적으로 다음 토큰을 예측하도록 학습시켜야 한다. 이때 학습을 위한 데이터는 GPT 계열 모델의 입력 형식에 맞추기 위한 특별한 토큰이나 형식을 포함하지 않은 순수 텍스트 시퀀스로 제공되어야 한다. 단, 하나의 문서가 매우 길 경우 잘라서 여러 시퀀스로 구성하고, 문서 간 구분을 위해 개행이나 구분자 토큰을 추가하는 방식을 사용하는 것은 가능하다. 또한 Unsloth는 주어진 말뭉치를 자동으로 샘플링하고 메모리 내에 패키징하여 효율적으로 학습을 진행하므로, 연구자는 학습률, 배치 크기 등의 필수적인 하이퍼파라미터만 설정하는 것으로 학습을 진행할 수 있다.

2.2.3 SFT 학습절차 구성

다음으로 SFT 단계에서는 Unsloth의 SFTTrainer를 이용하여 학습을 진행한다. SFTTrainer를 이용하는 경우 데이터셋을 사용자 프롬프트와 모델 응답의 쌍 형식으로 제공하여, 모델이 지정된 프롬프트에 대해 원하는 형태의 응답을 생성하도록 학습시킨다. 우리의 SFT용 데이터셋은 기업 문서를 활용한 가상 시나리오의 질의응답 및 요약 과제들로 이루어져 있다. 이를테면, 데이터 항목 하나는 "입력: 다음은 한 제품 제안서의 내용이다 ... [제안서 내용] ... 질문: 이 제안의 주요 목표는 무엇인가? / 출력: [주요 목표에 대한 모범 답변]"과 같은 구조로 작성되었다. 이러한 프롬프트-응답 쌍을 여러 개 모아 JSON 혹은 TSV 형태로 정리한 뒤, Unsloth에 해당 데이터셋을 로드하여 학습을 진행하였다. Unsloth는 다양한 챗봇 템플릿 (예: Vicuna 형식 등)을 지원하므로, 우리는 모델 종류(Qwen2.5 vs Gemma2)에 맞는 템플릿을 선택하여 프롬프트를 구성하였다. 예컨대 Qwen2.5-Instruct와 같은 모델의 경우 시스템 메시지와 사용자/모델 역할 지정을 포함한 포맷을 사용하였고, Gemma2-Chat 모델의 경우 그에 맞는 대화 포맷을 적용하였다. SFT 단계에서는 모델이 이미 CPT를 거친 상태의 가중치를 이어서 사용하므로, 이전 단계에서 학습된 도메인 지식을 토대로, 주어진 지시에 어떻게 답해야 하는지를 추가로 배우게 된다.

2.2.4 파인튜닝 절차 종합

이러한 2단계 파인튜닝 절차는, 첫째 단계에서 지식 획득, 둘째 단계에서 지식 활용 방법 습득으로 비유할 수 있으며, 최종적으로

로는 두 단계를 모두 거친 모델이 기업 문서를 잘 이해하고 적절히 활용할 수 있을 것으로 기대된다. 필요하다면 SFT 이후에 추가 미세 최적화(예: 사용자 피드백에 기반한 보강 학습 등)를 고려할 수 있지만, 본 연구에서는 CPT와 SFT 두 단계만으로도 충분한 개념 검증이 이루어진다고 판단하였다.

3. 결론

본 연구에서는 기업 내부 문서의 유형별로 최적화된 LLM 파인튜닝 방법론을 제안하고 구현하였다. 요약하면, 일방향 지식형 문서에 대해서는 추가 사전학습(CPT)을 통해 모델의 도메인 지식 기반을 확장하였고, 양방향 상호작용형 문서에 대해서는 지도 미세조정(SFT)을 통해 모델의 문맥 이해 및 응답 생성 능력을 향상시켰다. 이러한 이중 접근법은 기업 내부의 방대한 텍스트 자원을 효과적으로 활용하면서도, 필요한 부분에만 전문적인 데이터 제작 노력을 투입함으로써 효율적인 도메인 적응을 실현한다는 의미를 가진다. 특히 Unsloth 프레임워크를 통한 경량화 학습을 적용함으로써, 비교적 제한된 연산 자원으로도 최신 거대 언어모델을 튜닝하여 자체 활용형 LLM을 구축할 수 있음을 보였다. 이는 곧 중소 규모의 조직이나 전산 자원이 한정된 팀에서도 맞춤형 LLM 도입이 현실적으로 가능함을 시사한다.

결론적으로, 본 연구의 방법론은 한정된 데이터와 자원으로도 기업 환경에 특화된 LLM을 구축하기 위한 실질적 로드맵을 제시한다는 점에서 의미가 있다. CPT와 SFT의 결합 전략은 내부 지식 학습과 응용 능력 향상을 모두 달성하여, 모델이 기업 맞춤형 지능을 갖추도록 한다. 향후 더욱 발전된 데이터 생성 기법, 평가 연구, 및 실서비스 적용을 통해 본 방법론의 가치를 검증하고 강화한다면, 기업들은 자사의 지식자산을 최대한 활용하는 맞춤형 AI 비서를 확보하게 될 것으로 기대된다. 이러한 노력이 축적되어 LLM의 도메인 적응에 관한 이해가 깊어지면, 궁극적으로는 다양한 산업 분야에서 프라이빗하고도 실무 활용이 가능한 언어 모델 활용이 보편화될 수 있을 것이다.

참고문헌

- [2] unsloth.ai, Unsloth Documentation [웹사이트], (2025.4.14.),
URL: <https://docs.unsloth.ai/>
- [2] huggingface.co, unsloth/Qwen2.5-7B-Instruct [웹사이트], (2025.4.14.),
URL: huggingface.co/unsloth/Qwen2.5-7B-Instruct
- [3] Tobias Kerner, Domain-Specific Pretraining of Language Models: A Comparative Study in the

Medical Field, (2024.7.19.),

URL: arxiv.org/html/2407.14076v1

- [4] nature.com, Fine-tuning large language models for domain adaptation: exploration of training strategies, scaling, model merging and synergistic capabilities [웹사이트], (2025.04.14.),
URL: www.nature.com/articles/s41524-025-01564-y