



دانشکده مهندسی کامپیوتر

تولید خودکار داده آزمون در فازهای قالب فایل

پایان نامه برای دریافت درجه کارشناسی ارشد در رشته مهندسی کامپیوتر
گرایش نرم افزار

دانشجو:

مرتضی ذاکری نصرآبادی

استاد راهنما:

دکتر سعید پارسا

شهریور ۹۷



مختصری در
تاریخ

تأییدیه هیأت داوران جلسه دفاع از پایان نامه

نام دانشکده: دانشکده مهندسی کامپیوتر

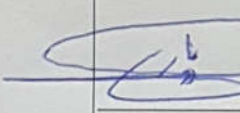
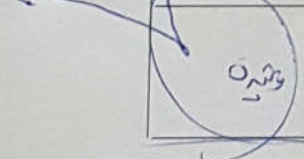
نام دانشجو: مرتضی ذاکری نصرآبادی

عنوان پایان نامه: تولید خودکار داده آزمون در فازهای قالب فایل

تاریخ دفاع: شهریور ۹۷

رشته: مهندسی کامپیوتر

گرایش: نرم افزار

ردیف	سمت	نام و نام خانوادگی	مرتبه دانشگاهی	دانشگاه / مؤسسه	امضا
۱	استاد راهنمای اول	دکتر سعید پارسا	دانشیار	دانشگاه علم و صنعت ایران	
۲	استاد راهنمای دوم	-	-	-	-
۳	استاد مشاور	-	-	-	-
۴	استاد مدعو داخلی	دکتر محمد عبداللہی ازگمی	دانشیار	دانشگاه علم و صنعت ایران	
۵	استاد مدعو خارجی	دکتر مجتبی وحیدی اصل	استادیار	دانشگاه شهید بهشتی	

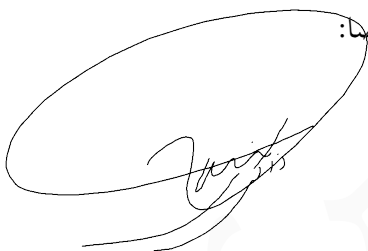
تأییدی صحت و اصالت نتایج

باسمه تعالی

اینجانب مرتضی ذاکری نصرآبادی به شماره دانشجویی ۹۵۷۲۳۰۸۸ دانشجوی رشته مهندسی کامپیوتر مقطع تحصیلی کارشناسی ارشد تأیید می‌نمایم که کلیه نتایج این پایان‌نامه حاصل کار اینجانب و بدون هرگونه دخل و تصرف است و موارد نسخه‌برداری شده از آثار دیگران را با ذکر کامل مشخصات منبع ذکر کرده‌ام. در صورت اثبات خلاف مندرجات فوق، به تشخیص دانشگاه مطابق با ضوابط و مقررات حاکم (قانون حمایت از حقوق مؤلفان و مصنفان و قانون ترجمه و تکثیر کتب و نشریات و آثار صوتی، ضوابط و مقررات آموزشی، پژوهشی و انضباطی ...) با اینجانب رفتار خواهد شد و حق هرگونه اعتراض درخصوص احقاق حقوق مکتسب و تشخیص و تعیین تخلف و مجازات را از خویش سلب می‌نمایم. در ضمن، مسئولیت هرگونه پاسخگویی به اشخاص اعم از حقیقی و حقوقی و مراجع ذیصلاح (اعم از اداری و قضایی) به عهده‌ی اینجانب خواهد بود و دانشگاه هیچ‌گونه مسئولیتی در این خصوص نخواهد داشت.

نام و نام خانوادگی: مرتضی ذاکری نصرآبادی

تاریخ و امضا:



11/28/2018
m-zakeri@live.com

مجوز بهره‌برداری از پایان‌نامه

بهره‌برداری از این پایان‌نامه در چهارچوب مقررات کتابخانه و با توجه به محدودیتی که توسط استاد راهنما

به شرح زیر تعیین می‌شود، بلامانع است:

☐ بهره‌برداری از این پایان‌نامه برای همگان بلامانع است.

☒ بهره‌برداری از این پایان‌نامه با اخذ مجوز از استاد راهنما، بلامانع است.

☐ بهره‌برداری از این پایان‌نامه تا تاریخ ممنوع است.

استاد راهنما: دکتر سعید پارسا

تاریخ:

امضاء:



تقديم

به دوستی که

می اندیشد،

می جوید،

می یابد

و

تغیر می دهد (:

قدردانی

افرادی در دوره مهیج و کوتاه کارشناسی ارشد و در روند پژوهش و نگارش این پایان نامه به بنده کمک کرده اند که برخورد لازم می دارم از زحمات ایشان سپاس گذاری نمایم. در ابتدا از استاد دانشمند و فرهیخته، آقای دکتر سعید پارسا که راهنمایی اینجانب را برعهده داشتند، صمیمانه قدردانی می کنم. بدون وجود حمایت ها و راهنمایی های ارزنده ایشان انجام این پایان نامه محقق نمی گردید. از داوران محترم پایان نامه آقایان دکتر محمد عبداللّهی ازگمی و دکتر مجتبی وحیدی اصل که با مطالعه و نقد سازنده خود، اسباب بهبود علمی و رفع نواقص این کار پژوهشی را فراهم نمودند، کمال تشکر را دارم. از همراهی ها و دلسوزی های همیشگی دو دوست و هم آزمایشگاهی بسیار عزیزم، آقای مهندس محسن امیریان و آقای مهندس سعید امیری، که بیش از ۶ سال با ایشان همکلاسی بودم، بی اندازه سپاس گزار هستم. از همکلاسی های بامعرفت و مهربانم در این دوره که همواره حامی و مشوق بنده بوده اند، لحظه های تکرار نشدنی را کنار ایشان تجربه کرده ام و متأسفانه مجال نام بردن از آنها را در این متن کوتاه ندارم، بی نهایت ممنون هستم. آشنایی با این عزیزان را دستاوردی ارزشمند در زندگی خود می دانم و برای شان بهترین ها را آرزو مندم. از دوستان گرامی، آقایان مهندس علی صابری و مهندس علی طاهری در آزمایشگاه رسانه دیجیتال دانشگاه صنعتی شریف، بابت انتقال مفاهیم اولیه یادگیری ژرف و آقای مهندس عرفان شرفزاده در آزمایشگاه سیستم های توزیع شده دانشکده مهندسی کامپیوتر، بابت پشتیبانی امور مربوط به پردازش ابری و ماشین های مجازی، تشکر می کنم. از آقای دکتر کارپتی برای وبلاگ و آموزش های مفید خود در زمینه یادگیری ژرف سپاس گزارم. از دیگر اساتید بزرگوaram در طول دوره کارشناسی ارشد از جمله آقایان دکتر محسن شریفی، دکتر بهروز مینایی و دکتر مهرداد آشتیانی (مدیر مرکز پردازش ابری دانشکده مهندسی کامپیوتر)، ممنون هستم. مهم تر از همه در پایان از خانواده عزیزم، به خصوص پدر و مادرم بابت حمایت های بی دریغ و همیشگی شان در همه مراحل زندگیم، مهربانانه سپاس گذاری کرده، دست ایشان را می بوسم.

مرتضی ذاکری نصرآبادی

شهریور ۹۷



11/28/2018
m-zakeri@live.com

تولید خودکار داده آزمون در فازهای قالب فایل

چکیده

آزمون فازی (فازینگ) یک فن آزمون پویای نرم افزار است. در این فن با تولید و تزریق مکرر داده های آزمون بدشکل به نرم افزار تحت آزمون (SUT)، به دنبال یافتن خطاها و آسیب پذیری های احتمالی موجود در آن هستیم. برای نیل به این هدف، آزمون فازی نیازمند داده های آزمون متنوع است. مشکل اساسی، پیچیده بودن ساختار ورودی برنامه هایی است که فایل را به عنوان ورودی می پذیرند. بررسی ها نشان می دهد بسیاری از داده های آزمون تولیدی در این موارد، مسیرهای محدود و سطحی را می پیمایند؛ زیرا در همان مراحل اولیه به علت بدشکل، بودن توسط تجزیه گر SUT رد می شوند. استفاده از ساختار گرامری فایل ها برای تولید داده ها، منجر به افزایش پوشش کد می شود؛ اما، استخراج گرامر برای ساختار فایل، اغلب، دستی صورت می پذیرد که مستلزم صرف هزینه و زمان زیاد و مستعد خطای فراوان است. در این پایان نامه، روشی خودکار برای تولید داده آزمون به صورت ترکیبی ارائه می دهیم. در روش خود، از مدل های زبانی عصبی (NLMS) که با شبکه های عصبی مکرر (RNNs) ساخته می شوند، استفاده می کنیم. مدل های پیشنهادی با فنون یادگیری ژرف، قادر به یادگیری آماری ساختار فایل های پیچیده و سپس تولید داده های جدید متنی، به صورت مبتنی بر گرامر و داده های دودویی به صورت مبتنی بر جابه جایی، هستند. فاز (بدشکل سازی) داده های آزمون، نیز توسط دو الگوریتم فاز جدید، تحت عنوان الگوریتم های فاز عصبی، که از این مدل ها استفاده می کنند، صورت می پذیرد. از روش پیشنهادی خود، برای تولید داده و سپس آزمون فازی نرم افزار پیچیده MuPDF که فایل های قالب سند حمل پذیر (PDF) را به عنوان ورودی می پذیرد، استفاده نمودیم. برای آموزش مدل های مولد، یک پیکره بزرگ از فایل های PDF را گردآوری کردیم. آزمایش های ما نشان می دهد که داده های تولید شده با این روش، منجر به افزایش میزان پوشش کد اجرایی SUT و بهبود بیش از ۷ درصدی آن، در مقایسه با فازهای قالب فایل مشهور، مثل AFL، می شود. آزمایش ها همچنین، بیانگر دقت یادگیری بهتر NLM های ساده تر در قیاس با مدل پیچیده تر کدگذار-کدگشا و نیز شکست این مدل، در میزان پوشش کد SUT، حین آزمون فازی، هستند.

واژگان کلیدی: آزمون فازی، داده آزمون، پوشش کد، یادگیری ژرف، شبکه عصبی مکرر.

فهرست مطالب

ر	فهرست شکل‌ها
ژ	فهرست جدول‌ها
س	فهرست الگوریتم‌ها
ش	فهرست گونه‌نوشت‌ها
۱	فصل ۱: مقدمه
۱	۱-۱ پیش‌زمینه
۲	۲-۱ شرح مسئله
۳	۱-۲-۱ شهود اولیه
۷	۱-۲-۲ کارهای مرتبط
۸	۱-۲-۳ فرضیه‌ها و اهداف
۹	۱-۳ روش پیشنهادی و نوآوری‌ها
۱۰	۱-۴ اهمیت موضوع
۱۲	۱-۵ ساختار پایان‌نامه
۱۳	فصل ۲: مفاهیم اولیه
۱۴	۲-۱ آزمون نرم‌افزار
۱۴	۲-۱-۱ معیارهای پوشش
۱۷	۲-۱-۲ دیدگاه جعبه

۱۸	۲-۱-۳ اندازه‌گیری پوشش کد
۱۹	۲-۲ آزمون فازی
۱۹	۲-۲-۱ فازر
۲۱	۲-۲-۲ معماری فازرها
۲۲	۲-۲-۳ آسیب‌پذیری
۲۳	۳-۲ یادگیری ژرف
۲۴	۲-۳-۱ عصب
۲۴	۲-۳-۲ شبکه عصبی روبه‌جلو
۲۵	۲-۳-۳ آموزش شبکه روبه‌جلو
۲۹	۲-۳-۴ شبکه عصبی مکرر
۳۲	۴-۲ مدل زبانی
۳۳	۲-۴-۱ مدل زبانی عصبی
۳۳	۲-۴-۲ ارزیابی مدل زبانی
۳۵	۵-۲ خلاصه

فصل ۳: کارهای مرتبط ۳۷

۳۸	۳-۱ فازر AFL
۳۹	۳-۱-۱ معماری AFL
۴۱	۳-۱-۲ مشکلات AFL
۴۲	۳-۲ فازر AFL افزوده
۴۳	۳-۲-۱ مشکلات AFL افزوده
۴۵	۳-۳ یادگیری و فاز
۴۵	۳-۳-۱ آموزش مدل کدگذار-کدگشا
۴۷	۳-۳-۲ تولید داده‌های جدید
۴۸	۳-۳-۳ اعمال آزمون فازی
۵۰	۳-۳-۴ مشکلات روش یادگیری و فاز

۵۱	۴-۳ دیگر کارهای مرتبط
۵۲	۵-۳ خلاصه
۵۵	فصل ۴: روش پیشنهادی
۵۵	۱-۴ کلیات
۵۶	۱-۴-۱ تمایز داده‌های متنی و دودویی
۵۸	۱-۴-۲ تمایز داده و فراداده
۵۸	۱-۴-۳ مثال انگیزشی
۵۹	۲-۴ مدل
۶۱	۲-۴-۱ آموزش مدل
۶۶	۲-۴-۲ تولید داده‌های جدید
۶۷	۳-۴ الگوریتم‌های فاز عصبی
۶۸	۳-۴-۱ فاز عصبی داده
۶۹	۳-۴-۲ فاز عصبی فراداده
۷۰	۴-۴ پیاده‌سازی
۷۳	۵-۴ خلاصه
۷۵	فصل ۵: ارزیابی روش پیشنهادی
۷۵	۱-۵ مورد مطالعاتی
۷۶	۲-۵ معیارهای ارزیابی
۷۹	۳-۵ چیدمان آزمایش‌ها
۷۹	۳-۵-۱ مجموعه داده
۸۰	۳-۵-۲ پیش‌پردازش داده‌ها
۸۳	۳-۵-۳ انتخاب فایل‌های میزبان
۸۳	۳-۵-۴ پوشش کد مبنا
۸۵	۴-۵ آزمایش‌ها، یافته‌ها و مقایسه نتایج
۸۶	۴-۵-۱ سرگشتگی، خطا و دقت مدل‌ها

۸۷	۵-۴-۲ پوشش کد مدل‌های مولد
۹۱	۵-۴-۳ مقایسه با مدل کدگذار- کدگشا
۹۱	۵-۴-۴ مقایسه در دوره‌های مختلف
۹۴	۵-۴-۵ آزمون فازی عصبی
۹۷	۵-۵ خلاصه
۱۰۰	فصل ۶: نتیجه‌گیری و کارهای آتی
۱۰۰	۶-۱ نتیجه‌گیری
۱۰۲	۶-۱-۱ مزایا و معایب یادگیری ژرف
۱۰۳	۶-۱-۲ مزایا و معایب آزمون فازی
۱۰۴	۶-۲ نوآوری‌ها
۱۰۵	۶-۳ ملاحظات اعتبارسنجی
۱۰۶	۶-۴ کارهای آتی
۱۰۸	مراجع
۱۱۳	پیوست آ: ساختار فایل PDF
۱۱۳	آ-۱ ساختار فیزیکی
۱۱۶	آ-۲ ساختار منطقی
۱۱۷	آ-۳ بروزرسانی فایل PDF
۱۱۸	پیوست ب: فازر IUST DeepFuzz
۱۱۸	ب-۱ پیش‌نیازهای نرم‌افزاری و سخت‌افزاری
۱۱۸	ب-۲ ساختار سطح بالای پروژه
۱۲۰	ب-۳ پیکربندی پروژه
۱۲۱	ب-۴ پیاده‌سازی‌ها
۱۲۱	ب-۴-۱ تعریف مدل‌ها
۱۲۳	ب-۴-۲ آموزش و تولید داده

ب-۴- پیاده‌سازی آزمون فازی ۱۲۳

ب-۵- ملاحظات عملی ۱۲۴

واژه‌نامه فارسی به انگلیسی ۱۲۵

واژه‌نامه انگلیسی به فارسی ۱۲۹

مستطبی
خاکری

فهرست شکل‌ها

۱۷	۲-۱ دیدگاه جعبه در آزمون نرم‌افزار
۲۰	۲-۲ روندنمای فرایند آزمون فازی
۲۵	۲-۳ ساختمان داخلی یک عصب
۲۶	۲-۴ گراف محاسباتی شبکه عصبی روبه‌جلو
۳۰	۲-۵ گراف محاسباتی شبکه عصبی مکرر
۳۱	۲-۶ انواع مختلف شبکه عصبی مکرر بر اساس طول توالی ورودی و خروجی
۳۴	۲-۷ مدل زبانی عصبی مکرر
۳۸	۳-۱ محیط زمان اجرای AFL
۳۹	۳-۲ مؤلفه‌های فازر AFL و ارتباط آنها
۴۶	۳-۳ مدل کدگذار-کدگشا
۴۷	۳-۴ مدل RNN کدگذار-کدگشا برای تولید اشیای داده‌ای PDF
۵۷	۴-۱ روندنمای روش پیشنهادی در حالت کلی
۶۰	۴-۲ مراحل یادگیری ساختار فایل، تولید و فاز داده آزمون برای فایل PDF
۶۲	۴-۳ گراف محاسباتی مدل‌های عصبی ژرف جدول ۴-۱
۶۵	۴-۴ مثالی از نحوه آموزش مدل‌های چند به یک
۷۴	۴-۵ معماری فازر قالب فایل پیشنهادی و ارتباط بین مؤلفه‌های مختلف آن
۸۴	۵-۱ پوشش کد برای هریک از میزبان‌ها و پوشش کدهای مبنا
۸۶	۵-۲ نمودار تغییرات خطای مدل‌های ۲ و laf

- ۵-۳ نمودار تغییرات پوشش کد مدل‌های مختلف برحسب تنوع. ۸۹
- ۵-۵ نمودار تغییرات پوشش کد برحسب دوره برای مدل‌های ۲ و laf. ۹۳
- ۵-۶ نمونه‌ای از داده‌های آزمون متنوع تولید شده توسط الگوریتم فاز عصبی داده در فرایند آزمون
فازی قالب فایل PDF. ۹۸
- آ-۱ یک فایل PDF باز شده در یک ویراستار متنی شامل عبارت Hello World ۱۱۴
- آ-۲ یک شیء PDF حاوی جریان دودویی ۱۱۶
- آ-۳ نمایش درختی ساختار منطقی فایل PDF ۱۱۶
- آ-۴ جدول ارجاع متقابل بروزرسانی شده یک فایل PDF ۱۱۷

فهرست جدول‌ها

- ۳-۱ مقایسه کارهای مرتبط ۵۳
- ۴-۱ مدل‌های یادگیری ژرف طراحی شده، پارامترها و ابرپارامترهای آنها. ۶۱
- ۵-۱ پارامترهای مؤثر در تولید خودکار داده آزمون با روش‌های یادگیری ماشینی و مقادیر قابل آزمایش برای آنها ۷۸
- ۵-۲ مشخصه‌ها و زمان آموزش مدل‌های جدول ۴-۱ در فرایند آموزش. ۷۹
- ۵-۳ نحوه و درصدهای تقسیم مجموعه داده به مجموعه‌های آموزش، ارزیابی و آزمون ۸۲
- ۵-۴ سرگشتگی، دقت و خطای مدل‌های مختلف روی مجموعه‌های آموزش و ارزیابی ۸۶
- ۵-۵ مقادیر مجاز و مقادیر داده شده به ثوابت و ورودی‌های الگوریتم‌های فازی عصبی داده و فاز عصبی فرا داده در هنگام آزمون فازی قالب فایل PDF ۹۵
- ۵-۶ نتایج پوشش کد حاصل از آزمون فازی نرم‌افزار MuPDF با الگوریتم‌های تولید داده آزمون مختلف ۹۵
- ۵-۷ میزان بهبود پوشش کد ابزار MuPDF توسط الگوریتم‌های روش پیشنهادی در مقایسه با کارهای مرتبط. ۹۵
- ب-۱ بسته‌های نرم‌افزاری مورد نیاز جهت اجرای صحیح برنامه IUST DeepFuzz. ۱۱۹

فهرست الگوریتم‌ها

۴۰		Basic-Random Fuzzing ۱-۳
۴۱		AFL Fuzzing ۲-۳
۴۴		Augmented-AFL Fuzzing ۳-۳
۴۹		SampleFuzz ۴-۳
۶۴		ProduceTrainingSamples ۱-۴
۷۱		DataNeuralFuzz ۲-۴
۷۲		MetadataNeuralFuzz ۳-۴

فهرست کوتاه‌نوشت‌ها

A

ASCII American Standard Code for Information Interchange
ASE Automated Software Engineering

B

BPTT Back Propagation Through Time

C

CAN Creative Adversarial Networks
CFG Control Flow Graph
CLI Command Line Interface
CNN Convolutional Neural Network

D

DAG Directed Acyclic Graph
DFG Data Flow Graph

G

GAN Generative Adversarial Network

GRU Gated Recurrent Units

GUI Graphical User Interface

I

ISRT Internet Security Threat Report

L

LM Language Model

LSTM Long-Short Term Memory

M

MOU Multiple Objects Update

N

NLM Neural Language Model

NLP Natural Language Processing

P

PDF Portable Document Format

R

RNN Recurrent Neural Network

S

SDL Security Development Lifecycle

SOU Single Object Update

SUT Software Under Test

فصل ۱

مقدمه

«من می‌گویم، امنیت، بالاترین اولویت ماست؛ زیرا برای همه چیزهای هیجان‌انگیزی که شما قادر به انجام دادن آن با کامپیوترها هستید – سازمان‌دهی زندگی‌تان، در ارتباط ماندن با دیگران، خلاق بودن – اگر ما مسائل امنیتی را حل نکنیم، مردم از همه این‌ها عقب خواهند ماند.»

❖ بیل گیتس

۱-۱ پیش‌زمینه

در حوزه مهندسی نرم‌افزار خودکار (¹ASE)، یکی از زمینه‌های مورد مطالعه و پژوهش، خودکارسازی فرایند آزمون نرم‌افزار، به‌عنوان یکی از مراحل مهم توسعه و ساخت یک سیستم نرم‌افزاری است. به‌طور کلی هدف از خودکارسازی، کاهش هزینه و زمان و افزایش دقت در اجرای یک فرایند است. آزمون فازی² یکی از فنون آزمون خودکار نرم‌افزار است. آزمون فازی در یافتن خطا³ها و آسیب‌پذیری⁴ها در نرم‌افزارهای دنیای واقعی مانند مرورگرهای وب، ویرایش‌گرهای متن، پخش‌کننده‌های چندرسانه‌ای و غیره، بسیار مؤثر واقع شده

¹Automated Software Engineering

²Fuzz Testing

³Fault

⁴Vulnerability

است [۱، ۲]. در این فن ورودی‌هایی بدشکل^۱ توسط یک برنامه دیگر، یعنی با روش خودکار، تولید شده و به نرم‌افزار تحت آزمون (SUT^۲) تزریق می‌شود. SUT در عین حال، به امید یافتن خطا بر اثر پردازش ورودی تزریق شده، پایش^۳ می‌شود. ورودی تولید شده که به برنامه داده می‌شود، نقش داده آزمون^۴ را داشته و عامل اصلی نمایان‌سازی خطا(های) احتمالی موجود در برنامه با بردن آن به یک حالت خرابی^۵ است. به‌همین علت مهم‌ترین مرحله در فرایند آزمون فازی را می‌توان تولید خودکار داده‌های آزمون دانست، به‌نحوی که بیشترین خطاها، ایرادها و آسیب‌پذیری‌ها شناسایی گردند.

۱-۲ شرح مسئله

راهکارهای مطرح در فن آزمون فازی [۳-۶]، برای شناسایی خطاها و آسیب‌پذیری‌ها نیازمند تولید تعداد زیادی داده آزمون هستند. در نرم‌افزارهایی با ساختار ورودی ساده، تولید داده آزمون نیز ساده است. برای مثال می‌توان با روش تصادفی این کار را انجام داد. اما در نرم‌افزارهایی با ساختار ورودی پیچیده، مانند فایل با قالب مشخص تولید داده آزمون متنوع که بتواند مسیرهای اجرایی بیشتری را پوشش دهد، کار آسانی نیست. تعداد و عمق مسیرهای اجرایی در یک برنامه با ساختار ورودی پیچیده به مراتب بیشتر از یک برنامه با ساختار ورودی ساده است. بررسی‌ها نشان می‌دهد بسیاری از داده‌های آزمون تولید شده برای چنین نرم‌افزارهایی، مسیرهای یکسان و سطحی (کم عمق) را می‌پیمایند [۷] و در مجموع، آزمون‌های فازی معمول پوشش کد ضعیفی دارند [۸]. درصد بالایی از داده‌های آزمون ساخته شده به‌صورت تصادفی، از لحاظ ساختاری کاملاً نامعتبر هستند و در همان مراحل اولیه بررسی صحت فایل، به‌وسیله تجزیه‌گر^۶ ورودی برنامه هدف، رد می‌شوند [۷، ۹]. در چنین شرایطی، قادر به نفوذ به عمق برنامه، کشف و آزمایش مسیرهای جدید نخواهیم بود. در واقع این نوع ورودی‌ها به نوعی تکراری و هدر رفته محسوب می‌گردند.

برای حل مسائل بالا، داده آزمون را با استفاده از قالب یا گرامر ورودی تولید می‌کنند، روش‌هایی مثل [۱۰]. قالب یا گرامر اما به صورت دستی و از روی مستندات تهیه می‌شود که با توجه به پیچیده بودن ساختار آن، عملی زمان‌بر، پرهزینه و مستعد خطا است [۱۱]. همچنین مستندات ساختار ورودی همواره در دسترس آزمون‌گر نیست. با این اوصاف روش مذکور تا به امروز، یکی از مؤثرترین روش‌های آزمون و یافتن خطا در برنامه‌هایی

^۱Malformed

^۲Software Under Test

^۳Monitor

^۴Test Data

^۵Failure

^۶Parser

مانند مرورگرهای وب بوده، که ساختار ورودی آن فایل‌هایی با قالب‌های متنوع و پیچیده هستند [۱۱، ۱۲]. به همین جهت، ارایه روشی برای خودکارسازی تولید داده آزمون بر مبنای قالب ورودی ارزشمند و حائز اهمیت است. پیش از ارایه یک روش جدید در ادامه ابتدا مسئله را دقیق‌تر تبیین کرده و راه‌حل‌های قبلی و نارسایی‌های هریک از آنها را مطالعه می‌کنیم.

۱-۲-۱. شهود اولیه

برای روشن شدن مسئله و شناسایی مشکلات موجود در تولید داده آزمون، مسئله را به زیر مسائل کوچک‌تر شکسته و از زوایای گوناگون تشریح می‌کنیم. ساختار پیچیده ورودی، ساختار پیچیده کد و تمایز داده و فراداده^۱ سه زیر مسئله‌ای هستند که ما آنها را شناسایی کرده و در این بخش، مطرح می‌کنیم. در ادامه این پایان‌نامه تمرکز خود را بر روی حل این مسائل منعطف خواهیم کرد.

ساختار پیچیده ورودی

نخستین مورد حائز اهمیت ساختار ورودی برنامه است. در برنامه‌هایی با ورودی خط فرمان (CLI)^۲ ساختارها به نسبت ساده هستند. اما برنامه‌هایی با ورودی فایل، ساختار ورودی بسیار پیچیده‌تری دارند. در واقع بسته به کاربرد، آنها یک یا چندین قالب فایل تعریف شده را پشتیبانی می‌کنند. همچنین برای یک قالب فایل شناخته‌شده ممکن است چندین نرم‌افزار وجود داشته باشد. یعنی در حالت کلی یک ارتباط چندبه‌چند بین قالب فایل ورودی و نرم‌افزار وجود دارد. هنگامی که یک نرم‌افزار برای آزمون انتخاب می‌شود هریک از قالب‌های فایلی که پشتیبانی می‌کند بخشی از کد نرم‌افزار را اجرا خواهند کرد.

فایل PDF^۳ را می‌توان نمونه‌ای از یک ورودی پیچیده برای نرم‌افزارهای PDF خوان، مثل اغلب مرورگرهای وب، تلقی کرد. مجموعه اسناد توصیف کننده مشخصه‌ها^۴ی کامل قالب PDF بیش از ۱۳۰۰ صفحه است [۱۱]. جزئیات ساختار این قالب را در پیوست آ بیان کرده‌ایم. در ساختارهای پیچیده هر بایت و در مواردی هر بیت نقش ویژه‌ای ایفا می‌کند که تولید تصادفی آنها تنوعی در پوشش کد برنامه ایجاد نمی‌کند؛ زیرا اغلب در دام کدهای کنترل استثنا^۵ می‌افتند. بنابراین داشتن یک درک حداقلی از ساختار در هنگام تولید داده جدید بسیار کمک کننده خواهد بود. چگونگی کسب این درک به‌صورت خودکار مسئله‌ای است که بایستی

^۱ Metadata

^۲ Command Line Interface

^۳ Portable Document Format

^۴ Specification

^۵ Exception Handling

حل شود.

ساختار پیچیده کد

پیچیده بودن ساختار ورودی، منجر به پیچیده شدن کد اجرایی و در نتیجه ممانعت از پوشش کد بالا در آزمون فازی خواهد شد. برای درک بهتر این مسئله برنامه ۱-۱ به زبان C را در نظر می‌گیریم. به خاطر سرعت بالای اجرا، بیشتر تجزیه‌گرهای قالب‌های پیچیده به این زبان نوشته می‌شوند. این برنامه یک فایل را از ورودی خوانده و براساس بایت‌های مشخصی در آدرس نسبی آن، مسیرهای معینی را اجرا می‌کند. چندین نکته قابل توجه در قطعه کد مذکور وجود دارد [۷]:

۱. **بایت‌های جادویی**^۱: بایت دوم و بایت اول ابتدا برای اعتبارسنجی ورودی با مقادیر ثابتی مقایسه می‌شوند. اگر نتیجه این مقایسه صحیح نباشد؛ ورودی درجا رد می‌شود. در سطر ۱۳ این مثال ابتدا آدرسی نسبی ۱ با مقدار 0xEF و سپس آدرس نسبی ۰ با مقدار 0xFD مقایسه می‌گردد. بایت‌های جادویی در قالب‌های فایل بسیاری وجود دارند. از جمله قالب فایل jpeg که در ابزار djpeg^۲ با همین روش، اعتبارسنجی می‌شود.

۲. **شرط‌های تودرتو**: در اجرای برنامه، هر مسیر اجرایی مهم است. هرچند رسیدن به برخی مسیرها ممکن است دشوارتر باشد یا حتی امکان‌پذیر نباشد [۱۳]. در این مثال برای رسیدن به خط ۱۸ کد بایستی همه شرایط موجود در خط ۱۷ برقرار باشد که به نوبه خود نیاز هست تا شرط موجود در خط ۱۵ نیز برقرار شد و به همین ترتیب. لذا داده آزمون تولیدی باید تا حد زیادی معتبر باشد تا بتواند به عمق مدنظر دسترسی پیدا کند.

۳. **نشان‌گرها**^۳: برای رسیدن به کد خط‌دار در سطر ۱۹ بایستی شرط سطر ۱۸ ارضا شود. این مقایسه با یک توالی از نشانه‌ها انجام می‌شود که آدرس نسبی شروع آن لزوماً ثابت نیست؛ البته در این مثال ثابت نشان داده شده است. در قالب‌های فایل‌هایی مانند png، jpeg و gif این قبیل نشان‌گرها دیده می‌شود.

۴. **آدرس‌های نسبی متغیر**: برای رسیدن به مسیر اجرایی سطر ۱۸ یک مقایسه بر مبنای آدرس‌های نسبی در سطر ۱۷ انجام می‌شود. آدرس نسبی به کار رفته در این مقایسه‌ها، از ورودی خوانده شده یا داخل برنامه محاسبه شده‌اند و بنابراین ممکن است که در هر بار اجرا متفاوت باشند. این امر برخلاف مورد بایت‌های جادویی است که آدرس نسبی ثابتی دارند.

^۱Magic Bytes

^۲<https://linux.die.net/man/1/djpeg>

^۳Markers

```

1  #include <stdio.h>
2  void main(int argc, char *argv[]){
3      unsigned char buffer[1024]; //Fixed size buffer
4      int fd, size, i, j;
5      /* Some initialization here */
6      if((fd = open(argv[1], O_RDONLY)) == -1)
7          exit(0);
8      fstat(fd, &s);
9      size = s.st_size;
10     if(size > 1024)
11         return -1;
12     read(fd, buffer, size);
13     if(buffer[1] == 0xEF && buffer[0] == 0xFD) //Complex
14         logic expression
15         printf("Magic bytes matched!\n");
16     else
17         EXIT_ERORR("Invalid input file\n");
18     if(buffer[i] == '%' && buffer[j] == '$' ){
19         if(strcmp(&buffer[15], "MAZE", 4) == 0) //Nested
20             condition
21             /* Codes contain bug here */
22         else{
23             /* *** Render file here (lines of code) *** */
24             close(fd);
25             return 0;
26         }
27     else{
28         EXIT_ERROR("Invalid bytes");
29         close(fd);
30         return 0;
31     }
32     close(fd);
33 }

```

برنامه ۱-۱: یک قطعه کد به عنوان نمونه‌ای از نرم‌افزار تحت آزمون در این پایان‌نامه، با ساختار تودرتو که چالش‌های پیچیدگی برنامه تحت آزمون و پوشش کد در آزمون فازی قالب فایل را نشان می‌دهد [۷] (با تغییر).

تمایز داده و فراداده

برنامه مبتنی بر ورودی فایل، به طور معمول دو گام مجزا را برای پردازش یک فایل طی می‌کند: گام اول تجزیه^۱ فایل و گام دوم پرداخت^۲ آن. در مرحله تجزیه، فایل در حافظه بارگذاری، مقادیر فیلدهای آن خوانده شده و تبدیل به داده‌ساختارهای داخل حافظه اصلی (مثل بافر، ساختمان یا رکورد، آرایه و غیره) می‌شود. در این مرحله چنانچه اشکال^۳ نحوی در ساختار فایل باشد (فایل از مشخصه‌های قالب خود پیروی نکند)، باید توسط تجزیه‌گر تشخیص داده شود وگرنه منجر به اشکال فساد حافظه^۴ و خرابی برنامه می‌شود. در مرحله پرداخت، برنامه روی اطلاعات خوانده شده از فایل پردازش لازم را انجام می‌دهد و خروجی تولید می‌کند (مثلاً نمایش یک تصویر روی صفحه نمایش یا اجرای یک ویدئو و غیره) [۱۴]. خطاهای این مرحله معمولاً جدی‌تر بوده و تشخیص آن نیز مشکل‌تر است، زیرا در عمق بیشتری از کد اجرایی رخ می‌دهند. جایی که داده‌های آزمون کمتری به آن دست پیدا می‌کنند.

با توجه به توضیح بالا، می‌توان یک فایل را حاوی دو دسته از مقادیر دانست: اول، مقادیری که مشخص کننده ساختار آن فایل هستند؛ برای مثال نام فیلدها. این مقادیر را فراداده یا دادگان (داده برای داده) می‌نامند. دوم، مقادیری که مشخص کننده اطلاعات هر فیلد هستند یا همان داده‌های فایل. مسئله نهفته در اینجا آن است که رویکرد آشکارسازی خطا برای هر کدام از این قسمت‌ها متفاوت خواهد بود؛ چراکه طبیعت خطاهای هر قسمت با یکدیگر متفاوت بوده و همان‌طور که گفته شد در مراحل مختلفی هم روی می‌دهند. برای آشکار کردن خطاهای تجزیه‌گر، لازم است تا فایل‌هایی تولید کنیم که بخش فراداده آن بدشکل شده‌اند در حالی که برای آشکار کردن خطاهای بخش پرداخت، بایستی فایل‌هایی تولید کنیم که از لحاظ نحوی معتبر بوده و بخش داده آن بدشکل شده باشند. شرط لازم هر دو نوع بدشکل‌سازی داشتن سازوکاری برای تشخیص داده و فراداده از یکدیگر، در هنگام تولید داده‌های آزمون است.

همان‌طور که گفتیم، رسیدن به کدهای مرحله پرداخت یک فایل (منظور تزریق داده آزمونی است که منجر به اجرای آن شود) سخت‌تر است. در برنامه ۱-۱، فرض شده است که به عنوان مثال پرداخت فایل در خط ۲۱ انجام می‌شود؛ یعنی، بعد از گذشتن از تمامی شرایط و بررسی‌های انجام شده توسط تجزیه‌گر و انتقال فیلدهای داخلی فایل به حافظه اصلی (فیلدهایی مثل فیلد size در برنامه مذکور). هر داده آزمونی که یکی از شرایط قبل از خط ۲۱ را نداشته باشد، رد شده و آن اجرا از برنامه به اجرای خط ۲۱ منتهی نمی‌گردد. برای آن که درصد خوبی از داده‌های آزمون تولید شده به اجرای خط ۲۱ منجر شوند، بایستی یک فایل تقریباً معتبر و

^۱Parse^۲Render^۳Error^۴Memory Corruption

پیروی کننده از قواعد قالب فایل مورد انتظار برنامه ۱-۱ را به عنوان داده آزمون تولید کرد. آنچه از شهود داده شده در این قسمت نتیجه می‌شود آن است که از یک برنامه قابل اطمینان و غیر قابل نفوذ، انتظار می‌رود که تحت هیچ عنوان بر اثر پردازش یک ورودی دچار خطا نشود. تنها زمانی می‌توان این ادعا را داشت که مطمئن شویم برنامه ورودی‌های به اندازه کافی متنوع را پردازش کرده و در هیچکدام از آنها دچار خطا نشده است. ورودی‌ها بایستی قادر به اجرای بخش‌های زیادی از کد برنامه باشند. حالت ایده‌آل اجرای تمام کد یک برنامه پیچیده است.

۱-۲-۲ کارهای مرتبط

تعدادی کار در ارتباط با استخراج خودکار گرامر فایل انجام شده‌اند. Bastani و همکاران [۱۵] الگوریتمی برای تولید یک گرامر مستقل از متن روی یک مجموعه از ورودی‌های نمونه داده شده ارایه کرده‌اند، که در نهایت برای تولید داده‌های جدید مورد نیاز آزمون فازی استفاده می‌شود. این الگوریتم یک مجموعه از مراحل تعمیم‌پذیری را با معرفی ساختارهای تکراری و متناوب برای عبارت‌های منظم به‌کار می‌بندد و غیر پایانه‌ها را برای گرامر مستقل از متن در هم ادغام می‌نماید که به‌نوبه خود یک گرامر یکنواخت از زبان ورودی به‌دست می‌دهد؛ اما، این روش برای قالب‌هایی مثل PDF که ساختار مسطح (غیر تو در تو) ولی در عین حال محتوای مختلفی از انواع و جفت‌های کلید-مقدار دارند، مناسب نیست [۱۱].

AUTOGRAM [۱۶] نیز به‌صورت غیر-احتمالاتی یک گرامر مستقل از متن را یاد می‌گیرد. یک مجموعه ورودی داده‌شده و به‌صورت پویا مشخص می‌شود که چگونه ورودی‌ها در برنامه پردازش می‌شوند. در واقع برنامه تحت آزمون با آلودگی پویا مشاهده می‌شود که حافظه را با قطعات ورودی که از آنها می‌آیند، برچسب‌گذاری می‌کند. بخش‌هایی از ورودی‌ها که توسط برنامه پردازش می‌شود، نهادهای نحوی در گرامر می‌شوند.

در پژوهش‌های اخیر تمایل زیادی به استفاده از شبکه‌های عصبی برای تحلیل و تولید برنامه‌ها به‌وجود آمده‌است. در سال ۲۰۱۷، Godefroid [۱۱] و همکاران روش جدیدی را برای تولید داده آزمون جهت استفاده در آزمون فازی بر مبنای مدل کدگذار-کدگشا^۱ [۱۷، ۱۸] ارایه کردند. در مقاله آنها، ساختار فایل PDF برای آزمون انتخاب شده است. ایده اصلی یادگیری یک مدل مولد روی مجموعه‌ای از ویژگی‌های اشیای داده‌ای PDF با داشتن مجموعه‌ای از نمونه‌های اولیه است. مدل کدگذار-کدگشا اجازه یادگیری متن با طول دلخواه را برای پیش‌بینی توالی بعدی کاراکترها، می‌دهد.

مدل استفاده شده توسط Godefroid و همکاران، مدل مبنایی وظایفی مانند ترجمه ماشینی یا تبدیل گفتار

¹Encoder-Decoder Model

به نوشتار است که یادگیری ساختار فایل را نمی‌توان در این وظایف گنجانده؛ زیرا، این مدل برای نگاشت دو توالی با دامنه‌های مختلف به کار گرفته می‌شود و این در حالی است که یادگیری ساختار فایل چنین وظیفه‌ای نیست. یعنی می‌توان از مدل‌های ساده‌تری مانند مدل زبانی نیز برای یادگیری ساختار فایل استفاده کرد. روش پیشنهادی آنها، تنها ساختارهای متنی فایل را مورد یادگیری قرار می‌دهد. این در حالی است که ساختار فایل‌های پیچیده هم متنی و هم دودویی هستند. افزون بر این، الگوریتم پیشنهادی آنها برای تولید داده آزمون نیز مشکلاتی دارد. از جمله اینکه ممکن است هیچ‌گاه پایان نیابد. در فصل ۲، ضمن تشریح کامل این روش، مشکلات آن را نیز به صورت کامل‌تری بیان می‌کنیم. همچنین در فصل ۳، دو فازر قالب فایل دیگر تحت عنوان AFL [۱۹] و AFL افزوده [۲۰] که با روش‌هایی غیر از یادگیری گرامر سعی در بهبود پوشش کد SUT در فرایند آزمون فازی را دارند، نیز بررسی می‌کنیم و مشکلات آنها را بیان خواهیم کرد.

در هیچ‌کدام از کارهای قبلی، مسئله مطرح شده در ارتباط با تمایز میان داده و فراداده در هنگام آزمون فازی دیده نشده است. به عبارت دیگر، این دیدگاه به آزمون فازی قالب فایل، دیدگاهی نو است و ارزش آزمایش شدن دارد. امکان استفاده از الگوریتم ارایه شده در روش [۱۱]، برای تمایز میان داده و فراداده وجود دارد اما برای حل مابقی مشکلات، یک روش جدید را در فصل ۴، پیشنهاد خواهیم داد.

۱-۲-۳ فرضیه‌ها و اهداف

هدف اصلی در پایان‌نامه پیش‌رو، ارائه روشی کارا جهت یافتن خطاها و آسیب‌پذیری‌ها در نرم‌افزارهایی مثل PDF خوان‌ها بوده که ورودی آنها فایل با ساختار مشخص و معمولاً پیچیده است. در این راستا تولید خودکار فایل‌های ورودی با هدف افزایش پوشش کد^۱ SUT از اهمیت ویژه‌ای برخوردار است. برای نیل بدین اهداف از فنون یادگیری ژرف^۲ در یادگیری و درک خودکار ساختار فایل و سپس تولید فایل‌های جدید، استفاده خواهیم کرد.

چون ایده استفاده از یادگیری ماشینی در آزمون فازی جدید است، این حوزه هنوز برای پژوهشگران ناشناخته بوده و بنابراین یکی دیگر از اهداف این پایان‌نامه شناسایی، تعریف و تفکیک پارامترهای حاکم در حوزه مذکور است. به نظر می‌رسد که فنون یادگیری ماشینی راه‌گشای حل مسائل شرح داده شده در بخش ۱-۲ باشد. به همین جهت فراهم آوردن چارچوبی استاندارد برای شکل‌دهی به کارهای آتی، مفید و مثمر ثمر خواهد بود. به طور خلاصه ما چندین فرضیه در این پایان‌نامه در نظر داریم، که تدوین سازوکارهایی برای رد یا تأیید صحت آنها، اهداف ما خواهند بود:

^۱Code Coverage

^۲Deep Learning

- استفاده از فنون یادگیری ژرف بالاًخص شبکه‌های عصبی مکرر ژرف، در یادگیری خودکار ساختار فایل، امکان‌پذیر و نتیجه‌بخش است.
- خودکارسازی کامل فرایند آزمون فازی مبتنی بر گرامر با ترکیب مدل یادگیری (مدل زبانی عصبی) و روش‌های فاز (بد-شکل‌سازی) ورودی، به خوبی میسر می‌شود.
- روش‌های ترکیبی تولید داده آزمون، یعنی روش تولید مبتنی بر گرامر به همراه روش تولید مبتنی بر جابه‌جایی، منجر به افزایش پوشش کد SUT می‌گردند.
- امکان کشف خطاها و آسیب‌پذیری‌های احتمالی موجود در SUT از طریق آزمون فازی با داده‌های آزمون تولید شده از طریق مدل‌های یادگیری ژرف، وجود دارد.

۳-۱ روش پیشنهادی و نوآوری‌ها

در این پایان‌نامه یک روش برای یادگیری خودکار ساختار فایل و سپس تولید داده‌های آزمون بر اساس آن ارائه می‌شود. برای یادگیری از مدل زبانی (LM^1) که یک مفهوم ابتدایی در پردازش زبان طبیعی (NLP^2)، است استفاده می‌کنیم. مدل زبانی را با استفاده از کلاس خاصی از شبکه عصبی ژرف^۳ موسوم به شبکه عصبی مکرر (RNN^4)، ایجاد می‌کنیم که در نتیجه به آن مدل زبانی عصبی (NLM^5) هم گفته می‌شود. روش ارائه شده در اینجا، همچنین، بخش‌های غیرمتنی را نیز در آزمون فازی لحاظ می‌کند، به هیچ قالب فایل خاصی وابستگی نداشته و به سبب پیاده‌سازی با زبان پایتون قابلیت اجرا بر روی هر ماشینی را دارد.

روش پیشنهادی در فصل ۴، در سه بخش کلی ارائه شده است. بخش اول به یادگیری ساختار فایل می‌پردازد (بخش ۴-۲)، بخش دوم روشی برای تولید و بدشکل‌سازی همزمان داده‌های آزمون ارائه می‌دهد (بخش ۴-۳) و در نهایت بخش سوم یک فازر کاملاً پیمانه‌ای را معرفی می‌کند که از آن برای آزمون فازی قالب فایل استفاده خواهد شد (بخش ۴-۴). در حالی که تمرکز اصلی بر روی نحوه تولید داده‌های آزمون است، اما برای انجام آزمون فازی به ابزارهای دیگری مانند تزریق کننده داده آزمون و نیز پایش SUT جهت ثبت خطاهای رخ داده نیاز است. نوآوری‌های اصلی روش پیشنهادی به طور خلاصه عبارتند از:

۱. یادگیری گرامر یا ساختار یک قالب فایل با استفاده از مدل‌های زبانی عصبی،

¹ Language Model

² Natural Language Processing

³ Deep Neural Network

⁴ Recurrent Neural Network

⁵ Neural Language Model

۲. تولید داده‌های آزمون متنی و دودویی همگام با بدشکل‌سازی آنها با استفاده از یک روش ترکیبی،

۳. ایجاد یک فازر قالب فایل و یک مجموعه داده آزمون برای آزمون فازی نرم‌افزارهای PDF خوان،

۴. و بررسی و شناسایی پارامترهای مؤثر در یادگیری ساختار فایل با استفاده از فنون یادگیری ژرف.

توضیح مبسوط‌تری از نوآوری‌های و دستاوردهای این پایان‌نامه در فصل ۶، ارائه شده است. در آن فصل همچنین مزایا و معایب فنون یادگیری ژرف در یادگیری ساختار فایل و نیز مزایا و معایب روش پیشنهادی بررسی و بیان شده‌اند.

۱-۴ اهمیت موضوع

مانند هر محصول دیگری، نرم‌افزار نیازمند آزمون^۱ و راستی‌آزمایی است. ماهیت غیرقابل لمس و پیچیدگی ذاتی نرم‌افزار سبب می‌شود تا فرایند آزمون آن نیز متفاوت، پیچیده و پرهزینه باشد. اما این دشواری‌ها از اهمیت موضوع آزمون نمی‌کاهد. خطاهای نرم‌افزاری در مواردی سبب خسارت‌های مالی و جانی جبران ناپذیری شده‌اند. راکت آریان ۲۵ اروپا در سال ۱۹۹۶، تنها ۳۷ ثانیه پس از پرتاب منفجر شد. علت آن وقوع خطا در تبدیل نوع یک عدد ممیز شناور به عدد صحیح بود [۱۳]. وجود خطا در ماشین پرتودرمانی Therac-25^۳، سبب کشته شدن دست‌کم سه انسان بر اثر تشعشع بیش از حد پرتو، در سال‌های ۱۹۸۵ تا ۱۹۸۷ شد. مثال‌های دیگری از این قبیل در [۱۳، ۲۱] آمده است.

در مواردی وجود خطا منجر به آسیب‌پذیری می‌شود که امکان سوء استفاده و دسترسی‌های غیرمجاز را به مهاجمان^۴ می‌دهد. باج‌افزار^۵ WannaCrypt^۶ که در نیمه اول سال ۲۰۱۷، بیش از ۲۳۰ هزار رایانه را در ۱۵۰ کشور جهان آلوده ساخت، از یک آسیب‌پذیری در هسته نسخه‌های قدیمی، سیستم عامل ویندوز شرکت مایکروسافت بهره‌برداری کرده بود. این باج‌افزار اطلاعات کاربر را رمزنگاری و برای رمزگشایی آن

^۱Test

^۲https://en.wikipedia.org/wiki/Ariane_5

^۳<https://en.wikipedia.org/wiki/Therac-25>

^۴Attacker

^۵Ransomware

^۶[https://docs.microsoft.com/en-us/windows/security/threat-protection/wannacrypt-ransomware-worm-targets-out-](https://docs.microsoft.com/en-us/windows/security/threat-protection/wannacrypt-ransomware-worm-targets-out-of-date-systems-wdsi)

[of-date-systems-wdsi](https://docs.microsoft.com/en-us/windows/security/threat-protection/wannacrypt-ransomware-worm-targets-out-of-date-systems-wdsi)

درخواست پرداخت هزینه می‌کرد. شرکت سیمنتک^۱ در گزارش ISRT^۲ خود در سال ۲۰۱۸ [۲۲]، افزایش ۶۰۰ درصدی حملات در اینترنت چیزها^۴ (اینترنت اشیاء) و افزایش تهدیدات در دیگر حوزه‌های سایبری از جمله تلفن همراه، را اعلام کرده است. در هر حال، کشف خطا و آسیب‌پذیری احتمالی ناشی از آن، در نرم‌افزارهایی که به طور گسترده توسط همگان مورد استفاده قرار می‌گیرند، مثل سیستم عامل‌ها، مرورگرهای وب، PDF خوان‌ها و غیره، بسیار حائز اهمیت است؛ زیرا، در صورت برطرف نشدن آن خطر وقوع حملاتی مشابه حملات بالا دور از انتظار نخواهد بود.

هنگامی که نرم‌افزارها بزرگ می‌شوند، آزمون دستی پاسخ‌گو نیست و خودکارسازی آزمون اهمیت می‌یابد. آزمون فازی همان‌طور که در ابتدای فصل بیان شد، به عنوان یک فن مؤثر آزمون نرم‌افزار در شناسایی خطاهای حافظه و آسیب‌پذیری‌ها شناخته شده است. برای مثال چرخه حیات امن نرم‌افزار (SDL^۵) شرکت مایکروسافت^۶، در مرحله درستی‌یابی^۷، استفاده از آزمون فازی را به عنوان یک روش استاندارد، اجباری می‌کند [۲۳]. آزمون فازی جعبه سفید^۸ (رجوع کنید به بخش ۲-۱-۲)، حدود یک سوم کل آسیب‌پذیری‌های شناخته شده در سیستم عامل ویندوز ۷ این شرکت را کشف کرده است [۱۰]. شرکت گوگل در سال ۲۰۱۲، اطلاعاتی راجع به ابزار ClusterFuzz خود منتشر کرد که از آن برای آزمون فازی پروژه‌های Chromium^۹ (شامل Chromium OS و مرورگر وب Chromium) استفاده می‌کند [۱۲]. این شرکت همچنین به افرادی که موفق به کشف آسیب‌پذیری در پروژه‌های نام‌برده شوند، جوایزی اهدا می‌کند.

تولید داده آزمون را بایستی مهم‌ترین مرحله در آزمون فازی دانست؛ چراکه داده‌هایی که نرم‌افزار با آنها آزمون می‌شود عامل اصلی اجرا شدن کدهای قسمت‌های مختلف SUT است و در صورتی که خطایی در آنها وجود داشته باشد، تنها از این طریق است که خود را نشان می‌دهد. البته باید بدین مسئله توجه کرد که اجرای کد خطا دار شرط لازم برای آشکارسازی خطا است ولی کافی نیست و روش تولید داده‌های آزمون، می‌بایست شرایط خاص بدشکل بودن را نیز محیا کند. تولید داده مبتنی بر گرامر، مؤثرترین روش آزمون برنامه‌هایی با ساختار ورودی پیچیده است [۲۴]. موارد بیان شده در این بخش، به‌خوبی اهمیت موضوع تولید خودکار داده آزمون در آزمون فازی و لزوم ارایه روش‌های جدید را توجیه کرده و انگیزه کافی را برای پژوهش در این

^۱ <https://www.symantec.com/>

^۲ Internet Security Threat Report

^۳ <https://www.symantec.com/security-center/threat-report>

^۴ Internet of Things

^۵ Security Development Lifecycle

^۶ <https://www.microsoft.com/en-us/sdl>

^۷ Verification

^۸ White Box

^۹ <https://www.chromium.org/>

زمینه ایجاد می کنند.

۱-۵ ساختار پایان نامه

این پایان نامه در شش فصل و دو پیوست تنظیم شده است و ساختار ادامه آن به قرار زیر است. در فصل ۲ ادبیات موضوع شامل آزمون نرم افزار، آزمون فازی و یادگیری ژرف را مطرح می کنیم. در این فصل ابتدا معیارهای سنجش کیفیت آزمون و چگونگی محاسبه آنها را توضیح داده، سپس به معرفی آزمون فازی، فرایند کلی و روش های تولید داده آزمون در آن می پردازیم. در بخش پایانی مباحث یادگیری ژرف را با تمرکز بر مفاهیم مرتبط با یادگیری ساختار فایل، عنوان خواهیم کرد.

در فصل ۳ به پیشینه پژوهش و بیان کارهای مرتبط در تولید خودکار داده آزمون و نقد و بررسی آنها می پردازیم. به طور خلاصه برخی راه حل های دیگران برای مسائل مطرح شده در بخش ۱-۲ را معرفی و سپس مشکلات آنها را بیان می کنیم. روش پیشنهادی در راستای حل این مسائل و ارزیابی ما در مقایسه با نتایج ارائه شده قبلی در این فصل خواهد بود.

در فصل ۴ روش پیشنهادی خود را برای تولید داده آزمون مطرح می کنیم. روش پیشنهادی در این فصل، همان طور که بدان اشاره شد، یک روش تولید مبتنی بر مدل های زبانی است که ما جابه جایی های تصادفی را نیز به آن اضافه کرده و روشی ترکیبی خلق نموده ایم. در همین فصل، ما دو الگوریتم جدید را برای فاز داده های آزمون معرفی می کنیم.

در فصل ۵ معیارهای ارزیابی روش پیشنهادی، چیدمان آزمایش ها و نتایج حاصل از اجرای آنها را ذکر خواهیم کرد. مورد مطالعاتی ما در آزمایش های این فصل نرم افزار MuPDF^۱ [۲۵] و قالب فایل PDF است که در ابتدای فصل آنها را مختصر معرفی خواهیم کرد.

در نهایت فصل ۶ را به بیان نتیجه گیری، یافته ها و نوآوری های پایان نامه، محدودیت های روش پیشنهادی و کارهای قابل انجام در آینده اختصاص داده ایم. همچنین در پیوست آ ساختار فایل PDF و در پیوست ب جزئیات پیاده سازی محصول نهایی پایان نامه را درج کرده ایم.

^۱<https://mupdf.com/>

فصل ۲

مفاهیم اولیه

«آزمون نرم افزار می تواند وجود خطاها را نشان دهد، اما هرگز نمی تواند نبود آنها را

تضمین کند.»

✠ ادسخر دایکسترا

در بخش اول این فصل تعاریف پایه ای آزمون نرم افزار و معیارهای سنجش آزمون، روش اندازه گیری و گزارش آنها را بیان می کنیم. در بخش دوم، آزمون فازی، فنون مختلف به کار گرفته شده در این آزمون، آزمون فازی قالب فایل، روش های تولید خودکار داده آزمون در فازر^۱ ها و معماری فازرها را توضیح می دهیم. بخش سوم و چهارم را به معرفی فنون یادگیری ژرف و استفاده از آنها در وظایف یادگیری مبتنی بر توالی^۲ تخصیص داده ایم. در پایان نیز خلاصه ای از مطالب مطرح شده در فصل را می آوریم. مفاهیم اولیه بسیار گسترده تر از آنچه در این فصل مطرح گردیده هستند و به همین دلیل در بخش پایانی خواننده را به منابع مناسبی ارجاع داده ایم. همچنین مطالعه کامل تری در باب مطالب این فصل در گزارش سمینار کارشناسی ارشد اینجانب [۲۶] انجام شده است.

^۱Fuzzer

^۲Sequence

۲-۱ آزمون نرم افزار

مجموعه فنون کشف و آشکارسازی خرابی های نرم افزار در مراحل مختلف توسعه آن را آزمون نرم افزار گویند. منظور از خرابی بروز رفتار (های) ناخواسته و خلاف مشخصه ها در یک نرم افزار یا قسمتی از آن است. خرابی خود حاصل یک خطا (نقص) ایستا در نرم افزار است. حالت داخلی نادرست برنامه را که ناشی از یک خطا است، اشکال می گویند [۱۳]. واژه های خطا، اشکال و خرابی از حوزه اتکاپذیری^۱ وارد آزمون نرم افزار شده اند [۲۱] و تعاریف یکسانی برای آنها وجود ندارد [۱۳، ۲۷].

آزمون نرم افزار منحصر به کد اجرایی برنامه نیست و شامل آزمون نیازمندی^۲ ها، آزمون مشخصه ها و طراحی^۳ نیز می گردد. در آزمون فازی اما مُنحصراً به کد اجرایی پرداخته می شود. لذا در ادامه این پایان نامه، منظور از آزمون، فقط آزمون مربوط به کد اجرایی برنامه خواهد بود. واژگان حاکم بر حوزه آزمون نرم افزار بسیار گسترده می باشد. در اینجا روی برخی از مفاهیم پایه ای و مرتبط با آزمون فازی و تولید داده آزمون تمرکز می کنیم.

۲-۱-۱ معیارهای پوشش

متأسفانه آزمون تنها قادر است وجود خرابی را نشان دهد و نبود آن را تضمین نمی کند. به بیان صوری، مسئله یافتن تمامی خرابی ها در یک برنامه تصمیم ناپذیر^۴ است [۱۳]. این موضوع سبب می شود به دنبال راه کارهایی برای اندازه گیری آزمون و تعیین میزان خوب بودن آن باشیم. در دهه های اخیر تمرکز بر روی معیارهای پوشش^۵، شناسایی، تفکیک و سنجش آنها بوده است. معیارهای پوشش بیشتر با هدف کمی سازی و اندازه گیری مقدار آزمون انجام شده مطرح شده اند؛ هرچند متقابلاً در تولید داده آزمون هم کاربرد دارند. چون با مسئله ای تصمیم ناپذیر مواجه هستیم، بایستی معیاری وجود داشته باشد که مشخص کند چه زمانی می توانیم آزمون را خاتمه دهیم و در این زمان آزمون تا چه حد خوب انجام شده است. آمان^۶ و آفوت^۷ [۱۳] معیارهای پوشش

¹Dependability

²Requirement

³Design

⁴Undecidable

⁵Coverage Criteria

⁶P.Ammann (<https://cs.gmu.edu/~pamman/>)

⁷J. Offutt (<https://cs.gmu.edu/~offutt/>)

را به صورت چهار معیار افراز فضای ورودی^۱، پوشش گراف^۲، پوشش منطق^۳ و پوشش ساختار نحوی^۴ مطرح کرده اند.

پوشش گراف در سطح کد اجرایی که به آن پوشش کد هم گفته می شود، شامل پوشش گراف جریان کنترل (CFG^۵) و گراف جریان داده (DFG^۶) برنامه می شود. پوشش منطق، مقداردهی و تعیین ارزش عبارات منطقی ظاهر شده در متن برنامه است. افراز فضای ورودی یعنی انتخاب از بین حالت های مختلف ترکیب ورودی ها و در نهایت پوشش ساختار نحوی، استفاده از قوانین گرامر برای اعتبارسنجی^۷ داده های ورودی یا تولید داده های جدید آزمون را شامل می شود.

معیارهای پوشش به دو روش مختلف قابل استفاده هستند. یک روش به این صورت است که مقادیر داده های آزمون منحصرراً برای برآوردن یک معیار داده شده، تولید گردند. این روش در برخی موارد بسیار دشوار است. به ویژه اگر ابزاری برای تولید خودکار داده آزمون نداشته باشیم یا ساختار ورودی پیچیده باشد. روش دوم تولید داده های آزمون به طور کاملاً مجزا و سپس اندازه گیری میزان پوشش معیار مربوطه است. برای روش اول یک برنامه شناسنده^۸ برآورده شدن معیار و برای روش دوم یک برنامه مولد^۹ داده آزمون نیاز است. در صنعت استفاده از روش دوم مرسوم تر است [۱۳].

در این پایان نامه، معیار دوم از معیارهای پوشش (پوشش کد) با روش دوم را مبنا قرار می دهیم. به این صورت که یک روش تولید خودکار داده آزمون و اندازه گیری میزان پوشش کد ارایه خواهیم داد. این فن آزمون، به خوبی عملی بوده و آزمون فازی معمول نیز بر اساس این روش است. پوشش کد خود در سطوح مختلفی مطرح است؛ پوشش دستور^{۱۰}، پوشش شاخه^{۱۱} و پوشش مسیر^{۱۲} سه سطح شناخته شده هستند. در پوشش دستور اجرای حداقل یک مرتبه هر دستور برنامه به عنوان نیازمندی تعریف و سپس اندازه گیری می شود. در پوشش شاخه اجرای حداقل یک مرتبه هر انشعاب برنامه به عنوان نیازمندی مطرح است و بالأخره در پوشش مسیر اجرای حداقل یک مرتبه هر مسیر اجرایی مد نظر است. بدیهی است که پوشش مسیر، پوشش شاخه و

¹Input Space Partitioning

²Graph Coverage

³Logic Coverage

⁴Syntax-Based Coverage

⁵Control Flow Graph

⁶Data Flow Graph

⁷Validation

⁸Recognizer Program

⁹Generator Program

¹⁰Statement Coverage

¹¹Branch Coverage

¹²Path Coverage

پوشش شاخه، پوشش دستور را شامل می شود. تعاریف و فهرست کاملی از سطوح پوشش کد و رابطه بین آنها در [۱۳] ذکر شده است.

شرکت مایکروسافت در محیط توسعه مجتمع ویژوال استادیو^۱، معیارهای پوشش خط^۲، پوشش بلوک پایه^۳ و پوشش خط جزئی^۴ را ارائه کرده است. پوشش خط همان پوشش دستور در کد سطح بالا (به جای کد اسمبلی/ماشین) است. پوشش بلوک پایه تعمیمی از پوشش دستور است که در آن یک توالی از دستورهای برنامه که شامل هیچ دستور پرشی در بین خود نیستند، یک دستور واحد در نظر گرفته می شوند. مزیت پوشش کد در سطح دستور و بلوک پایه^۵ سادگی اندازه گیری آن و عدم نیاز به وجود کد منبع^۶ برنامه است [۲۱]. پوشش خط جزئی منظور اجرا شدن برخی از قسمت های یک خط کد سطح بالا است. برای مثال در خطوطی با عبارات شرطی طولانی بر اثر ارزیابی مدار کوتاه^۷ ممکن است کلیه دستورات آن خط کد اجرا نگردد. در نتیجه پوشش خط جزئی معیاری متفاوت از پوشش خط است.

علاوه بر موارد ذکر شده معیارهای دیگری هم برای پوشش کد ارائه گردیده، از جمله پوشش دامنه^۸ که در آزمایشگاه تحقیقاتی مهندسی معکوس^۹ دانشگاه علم و صنعت ایران^{۱۰} توسعه داده شده است [۲۸]. پوشش دامنه، نوعی از معیار افزایش فضای حالت [۱۳] است که در آن با استناد به تئوری مجموعه ها، دامنه هایی برای هریک از ورودی های برنامه تعیین می شود. انتخاب ورودی از این دامنه ها در زمان آزمون، منجر به پوشش یک مسیر مشخص شده می گردد. از آنجایی که تنها یک بار اجرای یک مسیر، لزوماً خطای موجود در آن را آشکار نمی کند؛ در این روش می توان مسیر داده شده را به هر تعداد و هر بار با داده های آزمون متفاوت آزمون کرد. از این دیدگاه، پوشش دامنه یک روش تولید داده آزمون را ارائه می دهد.

در این پایان نامه، پوشش بلوک پایه و پوشش خط را برای گزارش میزان پوشش کد به کار می بریم. در بین این دو معیار، پوشش بلوک پایه در مقایسه با پوشش خط برنامه معیار مناسب تری است؛ زیرا، تعداد خطوط برنامه را می توان تغییر داد. برای مثال چندین دستور را در یک خط قرار داد یا اینکه یک دستور را در چند خط نوشت. تعداد بلوک پایه اما مستقل از نحوه آرایش و قرارگیری خطوط برنامه است و این معیار با تغییراتی مانند آنچه گفته شد، تغییر نمی کند؛ زیرا، در هر حالت تعداد بلوک های پایه یک برنامه و گراف جریان کنترل ثابت

^۱Visual Studio (visualstudio.microsoft.com)

^۲Line Coverage

^۳Basic Block Coverage

^۴Partial Line Coverage

^۵Basic Block

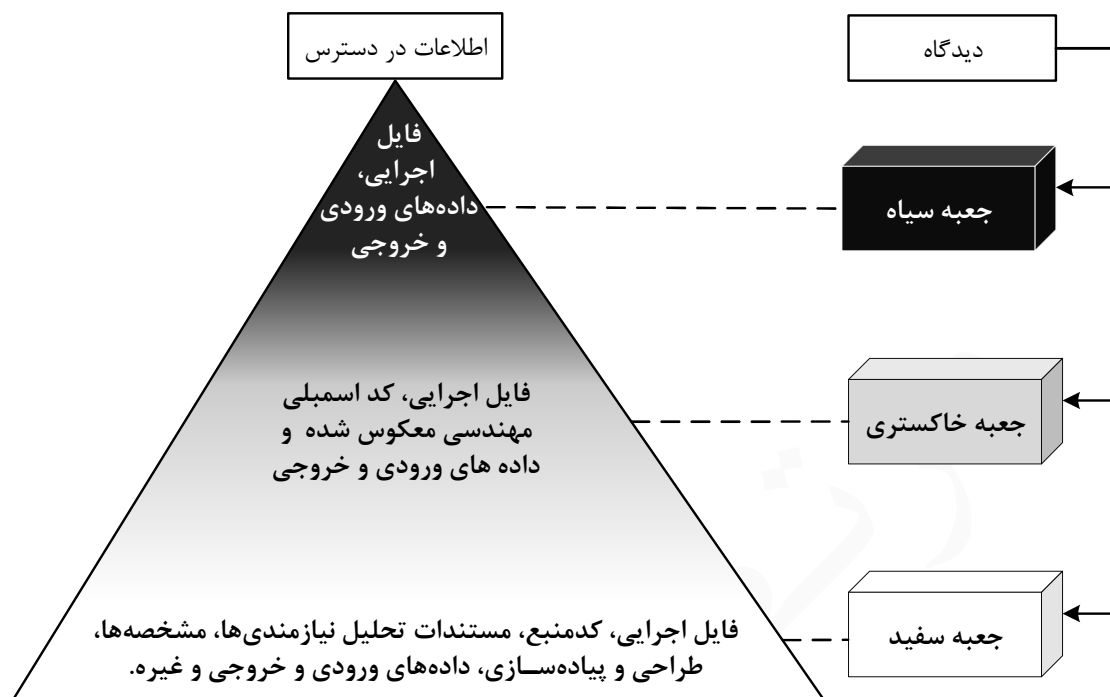
^۶Source Code

^۷Short-Circuit Evaluation

^۸Domain Coverage

^۹Reverse Engineering

^{۱۰}IUST Reverse Engineering Research Laboratory (<http://parsa.iust.ac.ir/reverse-engineering-lab/>)



شکل ۲-۱: طرح‌واره‌ای از دیدگاه جعبه بر اساس اطلاعات در دسترس به صورت یک مثلث آزمون [۲۶].

است. البته این تأکید صرفاً برای مواردی است که ممکن است آرایش خطوط برنامه در آزمون‌های مختلف تغییر کند، بدون اینکه تغییری در خود کد ایجاد شده باشد.

۲-۱-۲ دیدگاه جعبه

نحوه اندازه‌گیری معیارهای بحث شده، تنظیم و سنجش آزمون، منوط به اطلاعات در دسترس از SUT است. آزمون نرم‌افزار را براساس اطلاعات در دسترس از SUT در زمان آزمون، به سه دسته/ حالت جعبه سیاه^۱، جعبه سفید و جعبه خاکستری^۲ تقسیم‌بندی می‌کنند [۲۹، ۱۳]، که آن را دیدگاه جعبه به آزمون نرم‌افزار گوئیم. شکل ۲-۱ این تقسیم‌بندی را به صورت یک طرح‌واره مثلی نشان می‌دهد.

همان‌طور که در این طرح‌واره پیداست، در رأس مثلث، هیچ اطلاعی از ساختار داخلی SUT در اختیار آزمون‌گر نیست. در واقع با برنامه به مثابه یک جعبه سیاه که ورودی را دریافت و خروجی تولید می‌کند، رفتار می‌شود. این نوع آزمون معمولاً روی برنامه‌هایی که منتشر شده‌اند و توسط افرادی خارج از تیم توسعه و پشتیبانی محصول انجام می‌شود. هدف آن هم بیشتر کشف آسیب‌پذیری‌ها و ضعف‌های امنیتی است. در پایین

^۱ Black Box

^۲ Gray Box

مثلاً، دیدگاه جعبه سفید وجود دارد. در اینجا تقریباً تمامی اطلاعات SUT در دسترس است. دیدگاه جعبه خاکستری مابین دو دیدگاه قبلی قرار دارد. در این دیدگاه از فنون مهندسی معکوس برای استخراج اطلاعات از SUT استفاده می‌شود. هریک از این دیدگاه‌ها مزایا و معایب خاص خود را دارند [۳۰]. آزمون فازی هم که در ادامه مطرح می‌شود، قابل تقسیم به این سه دسته است. آزمون فازی جعبه خاکستری، برای مثال، بسیار ثمربخش بوده است [۱۹، ۲۰]. روش پیشنهادی در این پایان‌نامه برای تولید داده آزمون، قابل استفاده در هر سه دسته آزمون فازی است. برای ارزیابی روش پیشنهادی در این پایان‌نامه، در فصل ۵، اما از آزمون جعبه سفید، استفاده کرده‌ایم تا بتوانیم معیارهای پوشش کد و بخش‌های اجرا شده را به‌دقت اندازه‌گیری نماییم. از طرفی SUT انتخابی ما متن باز و رایگان بوده و در این حالت کد منبع آن در دسترس است.

۲-۱-۳ اندازه‌گیری پوشش کد

هدف از مطرح کردن دیدگاه جعبه در بخش ۲-۱-۲، اشاره به امکانات لازم برای اندازه‌گیری معیارهای پوشش از جمله پوشش کد بود. در آزمون جعبه سفید که کد منبع برنامه وجود دارد، مشاهده پوشش کد برنامه با اضافه کردن دستوراتی به متن کد به آسانی امکان پذیر است. این عمل را ابزارگذاری^۱ کد یا تجهیز برنامه به کد رهگیری می‌نامند [۱۴]. در محیط GNU/GCC (زبان‌های C و C++)، برای ابزارگذاری کافی است کد منبع برنامه را همراه با پرچم‌های `-fprofile-arcs` و `-fctest-coverage` کامپایل کنیم:

```
1 ~$ gcc -g -fprofile-arcs -fctest-coverage -o test test.c
2 ~$ ./test f1 f2
3 ~$ gcov test.c
4 ~$ more test.c.gcov
```

برنامه ۲-۱: ابزارگذاری کد در محیط GNU/GCC.

سطر اول برنامه ۲-۱ برنامه `test` را با پرچم‌های ذکر شده کامپایل می‌کند. سطر دوم برنامه را با دو آرگومان خط فرمان `f1` و `f2` اجرا می‌کند. سطر سوم پوشش کد اجرای اخیر برنامه را محاسبه و فایلی تحت عنوان `test.c.gcov` تولید می‌کند. در سطر چهارم این فایل برای مشاهده باز می‌شود. افزون بر این، پرچم‌هایی برای اخذ انواع سطوح پوشش کدی که بحث شد، از جمله پوشش شاخه، وجود دارد [۱۴].

محیط ویژوال استادیو ابزار `vsinstr` را برای ابزارگذاری کد برنامه و `VSPerfMon` را برای اندازه‌گیری پوشش کد زبان‌های پشتیبانی شده توسط این محیط، ارائه می‌دهد که البته در حضور کد منبع برنامه قابل استفاده

¹Instrumenting

هستند. هنگامی که کد منبع در اختیار نباشد کار اندکی دشوار می‌شود. در این موارد، بایستی از ابزارهای مهندسی معکوس و ابزارگذاری کد باینری استفاده کرد. ^۱ DynamoRIO، ^۲ Pin، ^۳ PaiMei، ^۴ QEMU و ^۵ [۳۳] ابزارهایی برای این منظور فراهم کرده‌اند.

۲-۲ آزمون فازی

آزمون فازی [۳-۶] همان‌طور که قبلاً هم گفتیم، فرایند ساده تولید و سپس تزریق یک ورودی ناخواسته (بدشکل شده یا نامتعارف) به SUT است. چنانچه برنامه بر اثر پردازش این ورودی ناخواسته دچار خرابی شود، حافظه برنامه مورد تحلیل قرار گرفته و خطای احتمالی موجود در کد آشکار می‌گردد. به دلیل اینکه SUT هنگام آزمون اجرا می‌شود، آزمون فازی، نوعی آزمون پویا است. چون SUT با تعداد ورودی‌های بسیار زیادی مورد آزمون قرار می‌گیرد، آزمون فازی را می‌توان نوعی آزمون فشار^۵ هم به‌شمار آورد. فرایند معمول آزمون فازی در شکل ۲-۲ نشان داده شده است.

۱-۲-۲ فازر

فازر ابزاری است که فرایند آزمون فازی را خودکار می‌کند. پیاده‌سازی هریک از پیمانه^۶‌های شکل ۲-۲ و تجمع آنها در کنار یکدیگر یک فازر را ایجاد می‌کند. تمرکز اصلی در یک فازر، نحوه تولید ورودی جدید به عنوان داده آزمون است به‌گونه‌ای که می‌توان آن را وجه تمایز اصلی فازرهای مختلف دانست. روش‌های تولید داده در فازرها قابل تفکیک به دو دسته کلی روش‌های مبتنی بر جابه‌جایی^۷ یا جهش و روش‌های مبتنی بر تولید^۸ هستند [۲۴].

در روش مبتنی بر جابه‌جایی، تعداد یک یا بیشتر داده ورودی معتبر^۹ به‌عنوان دانه اولیه^{۱۰} برای تولید داده‌های آزمون بیشتر به‌کار می‌رود. دانه اولیه جهش (تغییر ناگهانی) می‌یابد تا داده آزمون دیگری تولید شود.

^۱ <http://www.dynamorio.org/>

^۲ <https://software.intel.com/en-us/articles/pin-a-dynamic-binary-instrumentation-tool>

^۳ <https://github.com/OpenRCE/paimei>

^۴ <https://www.qemu.org/>

^۵ Stress Testing

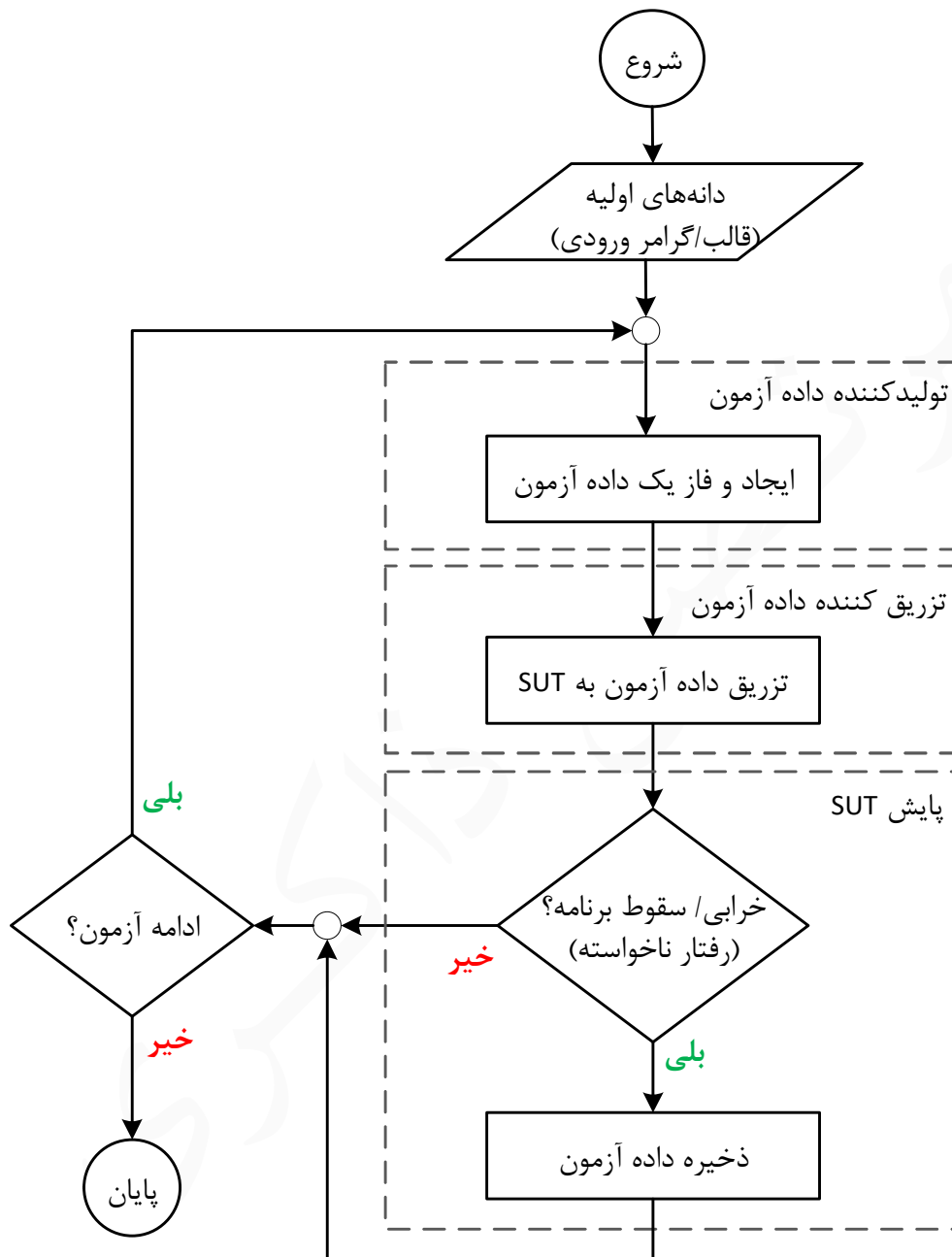
^۶ Module

^۷ Mutation Based

^۸ Generation Based

^۹ Valid

^{۱۰} Initial Seed



شکل ۲-۲: روندنمای فرایند آزمون فازی در حالت ساده که با اقتباس از مراجع مختلف ترسیم شده است. پیمانه‌های مورد نیاز برای خودکارسازی فرایند با مستطیل خط‌چین مشخص شده‌اند. شرایط ادامه آزمون می‌تواند بر عهده فرد آزمون‌گر قرار داده شود یا توسط خود فازر تعیین گردد.

جهش می‌تواند ساده باشد؛ شبیه معکوس کردن یک بیت، جایگزین کردن یک کاراکتر، یا پیچیده‌تر باشد؛ شبیه شناسایی و تکرار ساختارهای مشخص در داده، جایگزینی اعداد صحیح و حقیقی با مقادیر مرزی (بسیار بزرگ یا کوچک) و غیره. ساختن یک فازر مبتنی بر جابه‌جایی و تولید ورودی بد شکل و مخرب با آن، آسان است و نیاز به شناخت قبلی ساختار داده ورودی مورد جهش ندارد. عیب روش مبتنی بر جابه‌جایی اینجاست که این روش وابسته به تنوع ورودی‌های نمونه است. بدون وجود ورودی‌های مختلف با پیچیدگی کافی، این نوع فازر به پوشش کد بالایی دست نمی‌یابد [۱۲]. به عبارت دیگر دانه اولیه در فازر مبتنی بر جابه‌جایی بسیار حائز اهمیت است. AFL^۱ [۱۹]، FileFuzz^۲ [۲] و Radamsa^۳ [۱۲] نمونه‌هایی از فازرهای مبتنی بر جابه‌جایی هستند.

روش مبتنی بر تولید، داده‌های آزمون را به صورت کاملاً تصادفی یا از روی یک توصیف صوری مانند گرامر، قالب، یا مدل تولید می‌کند. در حالت دوم، از مشخصه‌های داده ورودی، برای ساخت یک مدل مولد^۴ استفاده می‌شود. این روش بیشتر روی قالب‌های داده‌ای که مستنداتی از مشخه‌های آنها در دسترس است، به کار می‌رود که در مقایسه با فازرهای مبتنی بر جابه‌جایی، معمولاً به پوشش کد بالاتری دست می‌یابد. با این حال، همان‌طور که در فصل ۱ هم اشاره شد، زمان و هزینه زیادی باید صرف شود تا مشخصه‌های یک قالب داده، کامل فهمیده و مدل خوبی از آن تهیه شود؛ زیرا، این کار تمام خودکار نیست [۲۴] و از طرفی، مستندات ساختار ورودی همواره در دسترس آزمون‌گر، نیستند. SPIKEfile^۵ [۲] و Peach^۶ نمونه‌ای از فازرهای مبتنی بر تولید هستند. روش‌های ترکیبی نیز وجود دارد که از ویژگی‌های هر دو روش بیان شده در تولید مورد آزمون کمک می‌گیرد. یک مثال از فازرهای ترکیبی LangFuzz [۳۵] است.

۲-۲-۲ معماری فازرها

یک معماری مرجع برای فازرها وجود ندارد؛ ولی می‌توان از روی فرایند ترسیم شده برای آزمون فازی (شکل ۲-۲)، پیمانه‌های اصلی لازم برای یک فازر را پیشنهاد کرد. پیمانه اول تولید کننده داده آزمون است که داده‌های ورودی لازم برای آزمون را تولید می‌کند. دیدیم که روش‌های مختلفی برای تولید داده آزمون وجود دارند. پیمانه دوم تزریق کننده داده آزمون است که وظیفه آن تحویل داده‌های تولید شده توسط تولیدکننده به

^۱ American Fuzzy Lop (<http://lcamtuf.coredump.cx/afl/>)

^۲ <http://www.fuzzing.org/>

^۳ <https://gitlab.com/akihe/radamsa>

^۴ Generative Model

^۵ <http://fuzzing.org/>

^۶ <https://www.peach.tech/>

SUT است. هر برنامه واسط مختص به خود را دارد. واسط می‌تواند CLI، گرافیکی (GUI^۱)، پروتکل شبکه یا فایل باشد. هدف این پایان‌نامه برنامه‌هایی هستند که ورودی آنها فایل است. معمولاً داده‌های ورودی به شکل فایل (مثل PDF) ساختار پیچیده‌تری نسبت به داده‌های CLI و پروتکل‌های شبکه دارند و تولید آنها سخت‌تر است. فازی که برای آزمون این برنامه‌ها توسعه داده می‌شود، فازر قالب فایل^۲ هم نامیده می‌شود [۲].

آخرین پیمانه مورد نیاز ابزار یا محیط پایش است که SUT را پایش می‌کند تا در صورت بروز خرابی ضمن ذخیره حالت برنامه، محل وقوع خرابی را بتواند با تحلیل حافظه برنامه مشخص کند. البته ابزارهای مستقل زیادی برای این منظور وجود دارند که می‌توان از آنها بهره گرفت. Application Verifier^۳ [۳۶] یک ابزار پایش قابل استفاده در سیستم عامل ویندوز است. ممکن است فازر، یک بازخورد از ابزار پایش برای تولید داده آزمون بعدی دریافت کند. این بازخورد عموماً حاوی اطلاعات مربوط به پوشش کد اجرای برنامه است. بر این اساس فازرها به دو دسته دارای حلقه بازخورد و بدون حلقه بازخورد تقسیم‌بندی می‌شوند [۲۴]. AFL [۱۹] یک فازر دارای حلقه بازخورد است. در این پایان‌نامه یک فازر ساده قالب فایل بدون حلقه بازخورد ارائه می‌دهیم. تمرکز بر روی نحوه تولید داده آزمون به روش مبتنی بر تولید / گرامر است.

۲-۲-۳ آسیب‌پذیری

آزمون فازی می‌تواند برای اهداف مختلفی استفاده شود، اما مهمترین هدف آن تحلیل نرم‌افزار به منظور کشف آسیب‌پذیری‌ها است. یعنی می‌توان آن را نوعی آزمون نفوذ^۴، آزمون قدرت‌مندی^۵ و آزمون امنیت^۶ هم تعریف کرد. تعاریف گوناگونی برای اصطلاح آسیب‌پذیری نرم‌افزار وجود دارد. آسیب‌پذیری را می‌توان اجتماع سه عنصر دانست: وجود یک خطا در نرم‌افزار، دسترسی مهاجم به خطا و قابلیت مهاجم برای بهره‌برداری^۷ از خطا [۲۴، ۳۷]. چنانچه پیش از این اشاره شد، خطا، نرم‌افزار را در یک حالت اشکال قرار می‌دهد که منجر به خرابی می‌شود. بعضی خطاها ممکن است تنها منجر به خرابی نرم‌افزار و بروز رفتار ناخواسته (مغایر با مشخصه‌ها) شوند. بنابراین صرف وجود یک خطا در نرم‌افزار آسیب‌پذیری محسوب نمی‌گردد [۲، ۱۴].

انواع مختلفی از خطاها وجود دارند که می‌توانند توسط فازرها کشف شوند؛ به‌ویژه آنهایی که از جنس نقض دسترسی به حافظه هستند. خطاهای فساد حافظه بدون شک رایج‌ترین و مؤثرترین روش بهره‌برداری

^۱Graphical User Interface

^۲File Format Fuzzer

^۳<https://docs.microsoft.com/en-us/windows-hardware/drivers/debugger/application-verifier>

^۴Penetration Testing

^۵Robustness Testing

^۶Security Testing

^۷Exploit

خرابکارانه از یک سیستم کامپیوتری محلی یا راه دور هستند. اگر حافظه بتواند به روشی خراب شود (یک آدرس بازگشت، یک اشاره گر پشته، یک اشاره گر تابع و غیره) اغلب اجرا می تواند به کد تهیه شده توسط مهاجم هدایت شود. بیشترین آسیب پذیری هایی که معمولاً توسط محققین حوزه امنیت کشف می شوند، مرتبط با میانگیر^۱ (یک ناحیه ثابت از حافظه اصلی در اختیار برنامه) هستند؛ از جمله سرریز میانگیر، سرریز پشته، سرریز عدد صحیح و غیره [۱]. در آزمون فازی به دنبال یافتن چنین خطاهایی از طریق تزریق ورودی بدشکل و خرابی حافظه SUT هستیم. مطالعه و طبقه بندی کاملی از انواع آسیب پذیری ها در [۲۶] انجام شده است.

۲-۳ یادگیری ژرف

یادگیری ژرف مجموعه ای از فنون یادگیری ماشینی است که در آن بردار ویژگی ها به صورت خودکار از داده های خام استخراج می گردد. ابزار اصلی و غالب در این فنون از یادگیری، شبکه های عصبی^۲ هستند. در حالت ساده شبکه های عصبی به مسئله یک تابع را نسبت داده و آن را مدل می کند. پارامترهای این تابع سپس در فرایند آموزش مدل، تخمین زده می شوند. در یادگیری ژرف تعداد این پارامترها صدها، هزارها و گاهی میلیون ها برابر افزایش می یابد که در نتیجه قدرت بازنمایی^۳ و حل مسئله مدل، افزایش خواهد یافت. شبکه عصبی در این حالت یک گراف با ژرفای بیش از یک یال، خواهد بود که به آن شبکه عصبی ژرف هم می گویند. شبکه عصبی ژرف بدین ترتیب قادر به استخراج و تمیز جزئی ترین ویژگی های مسئله و در نتیجه بهبود دقت حاصله، خواهد بود [۳۸].

وجود تعداد بسیاری پارامتر در مدل های یادگیری ژرف، سبب می شود که آموزش این مدل ها از لحاظ محاسباتی سنگین و زمان بر باشد. بسیاری از موانع تئوری و عملی آموزش این مدل ها به تازگی رفع شده اند. برای مسائل مختلف، مدل های متفاوتی با شبکه های عصبی ژرف طراحی و ساخته می شوند. شبکه عصبی روبه جلو^۴ ساده ترین نوع است. شبکه عصبی پیچشی (CNN)^۵ مناسب وظایف بینایی ماشین و شبکه عصبی مکرر (RNN) مناسب وظایف پردازش زبان طبیعی (NLP) (وظایف مبتنی بر توالی) ایجاد شده اند. وجه مشترک همه شبکه های نام برده، بلوک سازنده آن یعنی عصب^۶ است [۳۹].

¹ Buffer

² Neural Network

³ Representation

⁴ Feed-forward Neural Network

⁵ Convolutional Neural Network

⁶ Neuron

۲-۳-۱ عصب

بلوک محاسباتی پایه در بستر شبکه‌های عصبی، عصب یا واحد^۱ است. یک عصب تعدادی عدد حقیقی را به عنوان ورودی گرفته، محاسباتی روی آنها انجام داده و یک خروجی تولید می‌کند. محاسبه بدین صورت است که عصب جمع وزن‌دار ورودی‌هایی که به آن وارد می‌شوند را به همراه یک مقدار اضافی به عنوان بایاس^۲ حساب می‌کند. سپس برای جلوگیری از محدود شدن فضای توابع قابل تخمین به فضایی خطی، یک تابع غیرخطی، موسوم به تابع انگیزش^۳ (تابع فعالیت)، روی حاصل جمع اعمال می‌شود. توابع انگیزش متداول عبارتند از sigmoid و tanh. شکل ۲-۳ ساختمان داخلی یک عصب تنها با ورودی‌های x_1 ، x_2 و x_3 ، متقابلاً وزن‌های w_1 ، w_2 و w_3 و مقدار بایاس b را نشان می‌دهد. محاسبات این عصب مطابق رابطه‌های ۲-۱ و ۲-۲ است.

$$z = b + \sum_i w_i x_i \quad (۱-۲)$$

$$y = \sigma(z) \quad (۲-۲)$$

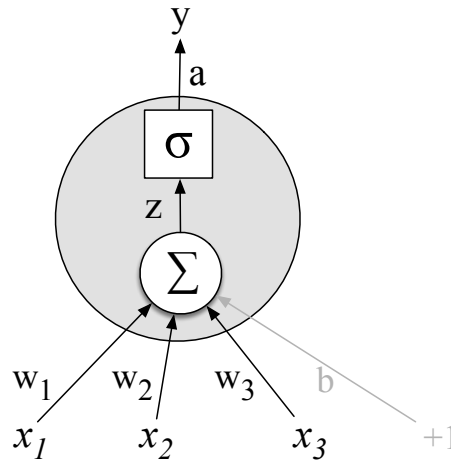
رابطه ۲-۱ را می‌توان به صورت ضرب داخلی بردارهای W و x هم نشان داد که نمایش مرسوم‌تری است [۳۹]:

$$z = b + W.x \quad (۳-۲)$$

۲-۳-۲ شبکه عصبی روبه جلو

شبکه روبه جلو را می‌توان با یک گراف جهت‌دار بدون دور (DAG^۴) نشان داد. هر گره یک عصب را نشان می‌دهد. گره‌هایی که مسیری به یکدیگر ندارند، یک لایه را تشکیل می‌دهند. ورودی هر لایه، خروجی عصب‌های لایه قبلی است. هر لایه بازنمایی یک مفهوم^۵ در قالب یک تابع است و مسئله ترکیبی از این مفاهیم

^۱Unit^۲Bias^۳Activation Function^۴Directed Acyclic Graph^۵Concept



شکل ۲-۳: ساختمان داخلی یک عصب با ورودی‌های x_1 ، x_2 و x_3 ، به ترتیب با وزن‌های w_1 ، w_2 و w_3 ، مقدار بایاس b و خروجی y [۳۹].

است. یک شبکه با دو لایه تابع $f(x) = W_2 \sigma(W_1 x)$ را پیاده می‌کند که W_1 و W_2 ماتریس پارامترهای تابع σ و تابع انگیزش است. $f(x) = W_3 \sigma(W_2 \sigma(W_1 x))$ معرف یک شبکه سه لایه خواهد بود. شکل ۲-۴ دو گراف محاسباتی با ژرفای ۲ و ۳ لایه (به ترتیب از چپ به راست) را نشان می‌دهد. لایه‌های میانی را لایه‌های پنهان^۱ هم می‌گویند. توجه شود که ورودی اول شبکه یک لایه محاسباتی نیست و بنابراین در شمارش لایه‌ها، شمرده نمی‌شود.

۲-۳-۳ آموزش شبکه روبه جلو

هدف از آموزش شبکه، تعیین مقادیر مناسب برای ضرایب و بایاس‌ها و سپس استفاده از شبکه برای انجام وظیفه^۲ مورد نظر است. داده‌هایی که در فرایند آموزش شبکه استفاده می‌گردد مجموعه آموزش^۳ نامیده می‌شوند. داده‌هایی که شبکه با آن مورد سنجش قرار می‌گیرد، مجموعه آزمون^۴ نامیده می‌شوند و مجموعه سوم از داده‌ها که در طول فرایند آموزش شبکه برای انتخاب بهترین راهکارهای یادگیری به کار گرفته می‌شوند، به نام مجموعه ارزیابی^۵ شناخته می‌شوند [۳۸].

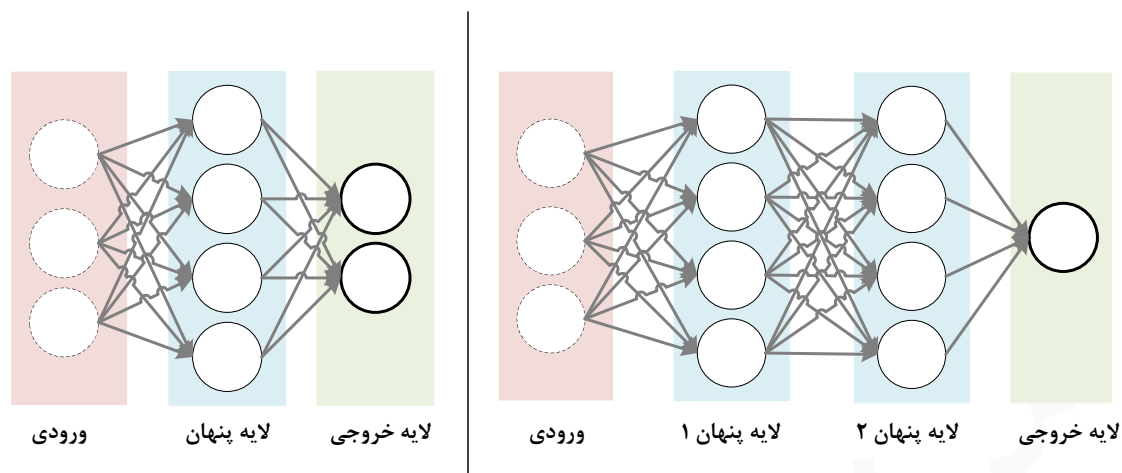
^۱Hidden Layer

^۲Task

^۳Training Set

^۴Test Set

^۵Validation Set



شکل ۲-۴: گراف محاسباتی برای دو شبکه عصبی روبه جلوی فرضی. چپ: یک شبکه عصبی دو لایه (یک لایه پنهان متشکل از چهار عصب و یک لایه خروجی متشکل از دو عصب) به همراه سه ورودی. راست: یک شبکه عصبی سه لایه (دو لایه پنهان هر کدام متشکل از چهار عصب و یک لایه خروجی متشکل از یک عصب) و به همراه سه ورودی.

شبکه عصبی در یک راهبرد بانظارت^۱ آموزش می بیند. در یادگیری بانظارت^۲ هر نمونه در مجموعه داده، علاوه بر بردار ویژگی ها با یک برچسب^۳ همراه است که نقش ناظر را ایفا می کند. طبقه بندی^۴ نمونه ای از وظایف یادگیری بانظارت است. در این وظیفه برای هر ورودی، کلاس خروجی آن در زمان آموزش مشخص است.

در فرایند آموزش ابتدا پارامترهای شبکه مقداردهی اولیه می شوند (معمولاً به صورت تصادفی). سپس به ازای هر ورودی از مجموعه آموزش، عصب های هر لایه خروجی خود را تولید می کنند تا زمانی که خروجی نهایی شبکه تولید شود. این مرحله را گذر جلو^۵ می گویند. از آنجایی که خروجی درست برای ورودی مربوطه در زمان آموزش در دسترس است (یادگیری بانظارت)؛ اختلاف مقدار محاسبه شده توسط شبکه ($\hat{y}$) و مقدار واقعی (y)، توسط تابعی موسوم به تابع خطا^۶ یا تابع هزینه^۷ یا تابع هدف محاسبه می شود [۴۰]. توابع مختلفی برای این منظور می توان تعریف کرد. تابع خطای مطلق میانگین^۸ (رابطه ۲-۴)، تابع خطای میانگین

¹Supervised

²Supervised Learning

³Label

⁴Classification

⁵Forward Pass

⁶Error Function

⁷Cost Function

⁸Mean Absolute Error

مربعات^۱ (رابطه ۲-۵) و تابع خطای آنتروپی متقاطع^۲ (رابطه ۲-۶) توابعی هستند که بسته به وظیفه مورد نظر استفاده می‌شوند [۴۱].

$$L_{MAE}(\hat{y}, y; W, b) = \frac{1}{n} \sum_{i=1}^n |\hat{y}_i - y_i| \quad (۲-۴)$$

$$L_{MSE}(\hat{y}, y; W, b) = \frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2 \quad (۲-۵)$$

$$L_{CE}(\hat{y}, y; W, b) = - \sum_{i=1}^n y_i \log \hat{y}_i \quad (۲-۶)$$

در روابط ۲-۴ تا ۲-۶، n تعداد نمونه‌ها یا تعداد مشاهده‌ها است. در این پایان‌نامه از تابع خطای آنتروپی متقاطع استفاده می‌کنیم که برای وظیفه‌های طبقه‌بندی با بیش از دو کلاس مناسب است [۴۰]. وقتی میزان خطا در گذر جلو مشخص شد، هدف تغییر میزان پارامترهای شبکه به نحوی است که خطا کمینه شود (رابطه ۲-۷):

$$f^* = \arg \min_{W, b} L(\hat{y}, y; W, b) \quad (۲-۷)$$

که f^* تابع یادگیری شده و مقادیر W و b پارامترهای شبکه هستند. سایر مقادیر را که در طول این فرایند یادگیری نمی‌شوند، ابر پارامتر^۳ می‌نامند. از جمله ابر پارامترها می‌توان به تعداد لایه‌ها و تعداد عصب‌ها در هر لایه اشاره کرد، که از آنها برای تعیین اندازه شبکه استفاده می‌گردد.

برای کمینه‌سازی خطا، مشتقات جزئی تابع خطا نسبت به هریک از پارامترها در لایه خروجی محاسبه می‌شود (گرادیان) و سپس این روند از طریق روال پس‌انتشار^۴ تا ابتدای شبکه پیش می‌رود و مقدار هر پارامتر در سوی کاهش گرادیان با ضربی موسوم به نرخ یادگیری^۵ بروزسانی می‌شود. نرخ یادگیری عدد کوچکی (در مقیاس صدم و هزارم) در نظر گرفته می‌شود که برای همگرایی آهسته و پیوسته شبکه و جلوگیری از

^۱ Mean Squared Error

^۲ Cross Entropy Error

^۳ Hyper-parameter

^۴ Backpropagation

^۵ Learning Rate

واگرا شدن آن لازم است. نرخ یادگیری یکی دیگر از ابر پارامترهای مورد نیاز آموزش شبکه‌های ژرف است. الگوریتم پسانتشار در واقع یک کاربرد بازگشتی از قاعده مشتق زنجیری در حساب دیفرانسیل است [۴۰]. در اینجا به بیان جزئیات الگوریتم‌های بهینه‌سازی نمی‌پردازیم؛ چارچوب‌های یادگیری ژرف انواع مختلفی از این الگوریتم‌ها را پیاده‌سازی و در اختیار قرار داده‌اند. توضیحات کاملی از این الگوریتم‌ها در [۲۸] آمده است.

شرایط توقف آموزش

فرایند آموزش شبکه می‌تواند تا بی‌نهایت ادامه داشته باشد. بنابراین بایستی یک قانون برای توقف آن تعریف شود. گزینه‌های زیادی وجود دارد که تحت چه شرایطی آموزش را خاتمه دهیم. همچنین ترکیب این شرایط نیز امکان‌پذیر است. برخی از این شرایط عبارتند از [۴۱]:

۱. هنگامی که تعداد مشخصی دوره^۱ طی شود (هر دوره برابر تعداد تکرارهایی است که برای مرور کل مجموعه آموزش لازم است).

۲. هنگامی که خطای مجموعه ارزیابی (برای تعداد مشخصی دوره متوالی) کاهش پیدا نکند.

۳. هنگامی که تغییرات میزان خطا (برای تعداد مشخصی دوره متوالی) کمتر از یک حد آستانه شود.

۴. هنگامی که یک زمان مشخص از فرایند آموزش سپری شود.

منظم‌سازی

مشکل رایجی که هنگام آموزش شبکه عصبی باید جلوگیری شود اثر بیش‌برازش^۲ است. بیش‌برازش یعنی با کاهش خطا روی مجموعه آموزش، خطای مجموعه‌های ارزیابی و آزمون به‌طور ناگهانی افزایش یابد [۳۸]. یک دلیل می‌تواند به‌اندازه کافی بزرگ نبودن اندازه مجموعه آموزش باشد. دلیل دیگر می‌تواند پیچیدگی بسیار بالای مدل باشد. پیچیدگی مدل با تعداد پارامترهای آن مشخص می‌شود. در پیچیدگی بالا ممکن است مدل، اختلال^۳‌های موجود در ورودی‌ها را نیز لحاظ کند و متناسب یا برانده آن شود. در نتیجه هنگام ارزیابی مدل روی داده‌های مجموعه آزمون، خطا بالا می‌رود [۴۱].

راهکارهای مختلفی برای مقابله با مسئله بیش‌برازش مدل پیشنهاد شده است. از این راهکارها تحت عنوان فنون منظم‌سازی^۴ یا تنظیم مدل یاد می‌شود. اغلب روش‌های شناخته شده، پارامترهای دارای مقادیر بالا را

^۱Epoch

^۲Overfitting

^۳Noise

^۴Regularization

که اثرات نوسانی شدیدی دارند، با تابعی جریمه می‌کنند [۴۰]. روش دیگر برای منظم‌سازی Dropout [۴۲] است که بسیار مؤثر و ساده بوده و در مدل‌های یادگیری ژرف معمولاً از آن استفاده می‌شود. در طول آموزش Dropout یک عصب را فقط با احتمال p (یک ابر پارامتر) فعال نگه می‌دارد و در غیر این صورت آن را صفر می‌کند. در نتیجه پیچیدگی مدل تنظیم می‌شود.

۲-۳-۴ شبکه عصبی مکرر

شبکه‌های روبه‌جلو در حل مسائلی که ترتیب ورودی در آنها مهم است، مثل ترجمه ماشینی و سایر وظیفه‌های NLP، نارسایی دارند. برای مثال در ترجمه ماشینی هر واژه به واژه‌های قبلی خود وابسته است. در این وظیفه‌ها ورودی به صورت یک توالی $x = \langle x^{(1)}, x^{(2)}, x^{(3)}, \dots, x^{(n)} \rangle$ است. غالب وظایف حوزه پردازش زبان طبیعی بدین صورت هستند.

شبکه‌های عصبی مکرر کلاسی از شبکه‌های عصبی هستند که به صورت یک گراف جهت‌دار دارای دور بیان می‌شوند. به عبارت دیگر ورودی هریک از لایه (های) پنهان یا لایه خروجی افزون‌بر خروجی لایه قبل، شامل ورودی‌ای از گام زمانی^۱ قبل به صورت بازخورد نیز می‌شود. شکل ۲-۵ یک RNN را نشان می‌دهد که در آن لایه پنهان از گام‌های زمانی قبلی بازخورد گرفته است. در هر گام زمانی t (از $t = 1$ تا $t = T$) یک بردار $x^{(t)}$ از توالی ورودی پردازش می‌شود. معادلات گذر جلو شبکه در t عبارتند از [۳۸]:

$$z^{(t)} = Ux^{(t)} + Wh^{(t-1)} + b \quad (۸-۲)$$

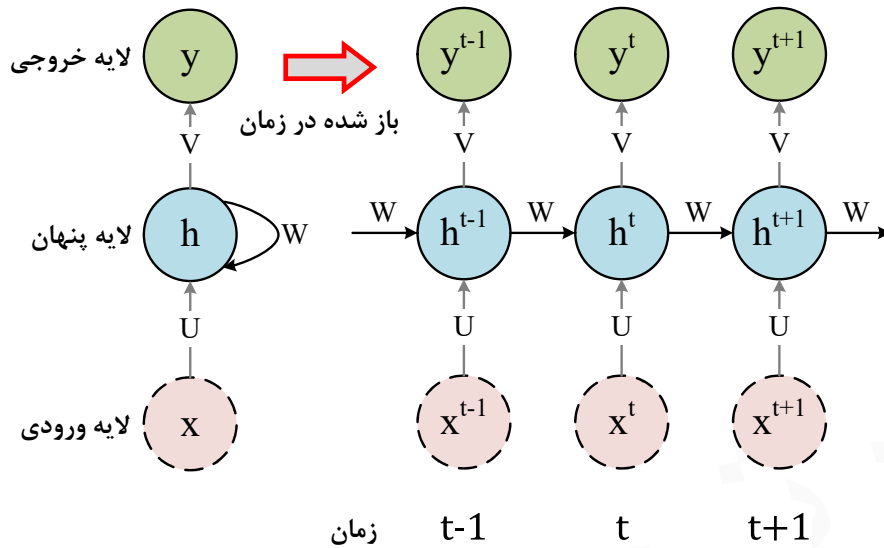
$$h^{(t)} = \sigma(z^{(t)}) \quad (۹-۲)$$

$$y^{(t)} = Vh^{(t)} + c \quad (۱۰-۲)$$

$$\hat{y}^{(t)} = \text{softmax}(y^{(t)}) \quad (۱۱-۲)$$

در روابط ۲-۸ تا ۲-۱۱، b و c بایاس و ماتریس‌های U ، V و W به ترتیب وزن یال‌های لایه ورودی

¹ Time Step



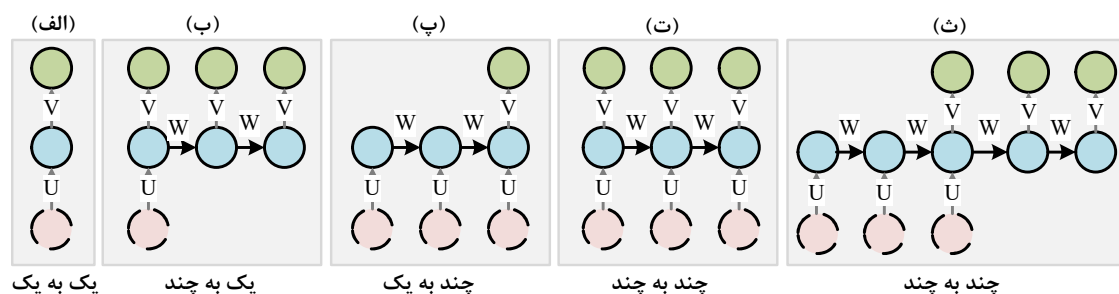
شکل ۲-۵: گراف محاسباتی مربوط به یک شبکه عصبی مکرر با یک لایه پنهان که یک توالی ورودی از مقادیر $x = \langle x^{(1)}, x^{(2)}, x^{(3)}, \dots, x^{(n)} \rangle$ را به یک توالی خروجی از مقادیر $y = \langle y^{(1)}, y^{(2)}, y^{(3)}, \dots, y^{(n)} \rangle$ نگاشت می‌کند. فرض شده است که خروجی y احتمال‌های نرمال نشده است، بنابراین خروجی نهایی شبکه یعنی $\hat{y}$ از اعمال تابع بیشینه هموار روی y حاصل می‌شود. چپ: شبکه عصبی مکرر به صورت یال بازگشتی (گراف جهت‌دار دارای دور). راست: همان شبکه به صورت باز شده در زمان، به نحوی که هر گره با برچسب مرحله زمانی مشخص شده است [۳۸].

به پنهان، پنهان به خروجی و پنهان به پنهان، پارامترهای شبکه هستند. در لایه خروجی به جای تابع انگیزش، تابع بیشینه هموار^۱ اعمال می‌شود. این تابع خروجی شبکه را به شکل یک توزیع احتمالی معتبر تبدیل می‌نماید و معمولاً در لایه خروجی مدل‌هایی که برای طبقه‌بندی استفاده می‌شوند، قرار می‌گیرد. تابع بیشینه هموار یک بردار k تایی از اعداد حقیقی را به عنوان ورودی دریافت نموده و یک بردار k تایی از مقادیر حقیقی در بازه $[0, 1]$ را به عنوان خروجی می‌دهد به طوری که جمع مؤلفه‌های آن برابر یک خواهد بود (رابطه ۲-۱۲):

$$\text{softmax}(x_i) = \frac{e^{x_i}}{\sum_{j=1}^k e^{x_j}} \quad \text{for } i=1 \text{ to } k \quad (2-12)$$

در شکل ۲-۵، RNN با یک لایه پنهان نشان داده شده است. اما می‌توان RNN ژرف با چندین لایه پنهان نیز داشت. همچنین طول توالی‌های ورودی و خروجی می‌تواند بسته به مسئله مورد نظر متفاوت باشد.

^۱Softmax



شکل ۲-۶: طرح‌واره‌ای از انواع حالت‌های مختلف شبکه عصبی مکرر براساس طول توالی ورودی و طول توالی خروجی. (الف): شبکه عصبی استاندارد، (ب): شبکه یک به چند، (پ): شبکه چند به یک، (ت) و (ث): شبکه‌های چند به چند [۴۰].

کارپتی^۱ [۴۰] RNN ها را از منظر طول توالی ورودی و طول توالی خروجی به چند دسته تقسیم‌بندی کرده است. شکل ۲-۶ این دسته‌بندی را نشان می‌دهد.

آموزش شبکه عصبی مکرر

الگوریتم پس‌انتشار برای آموزش RNN هم قابل استفاده است. در واقع RNN همان شبکه عصبی روبه‌جلو است که یک بازخورد از وزن‌های گام زمانی قبل دارد. در نتیجه خطا هنوز به‌صورت روی‌به‌عقب از گام زمانی t منتشر می‌شود. بسته به طول توالی ورودی و میزان منابع محاسباتی این انتشار تا گام زمانی $t = 1$ ادامه می‌یابد یا تحت محدودیت تعداد مشخصی گام^۲ متوقف می‌شود. نسخه‌ای از الگوریتم پس‌انتشار که می‌تواند به‌کل طول توالی ورودی اعمال شود، پس‌انتشار در زمان (BPTT^۳) نام دارد. محدودیت منابع محاسباتی ایجاب می‌نماید که طول توالی ورودی مقدار مشخصی قرار داده شود یا پس‌انتشار در تعداد محدودی گام زمانی متوقف شود [۳۸].

حافظه کوتاه‌مدت بلند

RNN در حالت استاندارد، برای توالی‌های طولانی که وابستگی بین ورودی‌هایی با فاصله زیاد وجود دارد، قادر به حفظ این وابستگی‌ها نیست؛ به بیان دیگر گرادیان در دوره‌های زمانی بلند مدت تمایل به ناپدید شدن^۴ یا انفجار^۵ (زیاد شدن بی‌اندازه) دارد. در حالت ناپدید شدن گرادیان، یادگیری مدل متوقف می‌شود. مدل

^۱ A. Karpathy (<http://karpathy.github.io/>)

^۲ Step

^۳ Back Propagation Through Time

^۴ Vanishing

^۵ Explosion

حافظه کوتاه مدت بلند ($LSTM^1$) برای مقابله با مسئله بالا و افزایش قدرت به خاطر سپاری RNN ایجاد شده است. هسته این مدل سلول حافظه C است که در آن برخلاف شبکه مکرر استاندارد که خروجی هر واحد تنها با اعمال یک تابع انگیزش ایجاد می شود، خروجی با محاسبات پیچیده تری تعیین می شود که به ویژه هدف آنها منظم سازی و حفظ قدرت گرادیان است. ایده اصلی در LSTM تخصیص مسیری مجزا علاوه بر یال بازخوردی حالت پنهان، برای جریان حافظه است.

جزئیات عملکرد LSTM در [۳۸] توضیح داده شده است. امروزه به طور پیش فرض از LSTM در پیاده سازی RNN استفاده می شود [۴۳]. البته راه کارهای جایگزین دیگری مشابه LSTM، نیز مطرح شده اند مانند GRU^2 [۱۸]، اما عملکرد آنها بهتر از LSTM نیست. مطالعه جامعی روی RNN و انواع آن برای یادگیری توالی ها در [۴۴] انجام شده است. کلیه مدل های یادگیری ژرف معرفی شده در این پایان نامه از LSTM در هسته خود استفاده می کنند.

۲-۴ مدل زبانی

در فصل ۱ به استفاده از LM در روش پیشنهادی خود، اشاره کردیم. LM یک مفهوم پایه در NLP است که امکان پیش بینی نشانه بعدی در یک توالی را فراهم می کند. به بیان دقیق تر LM عبارت است از یک توزیع احتمالی روی یک توالی از نشانه ها (اغلب واژه ها) که احتمال وقوع یک توالی داده شده را مشخص می کند. در نتیجه می توان بین چندین توالی داده شده برای مثال چند جمله، آن را که محتمل تر است، انتخاب کرد. LM برای توالی $x = \langle x^{(1)}, x^{(2)}, x^{(3)}, \dots, x^{(n)} \rangle$ به صورت زیر تعریف می شود [۴۵]:

$$p(x) = \prod_{t=1}^n p(x^{(t)} | x^{<t}) \quad (2-13)$$

در رابطه ۲-۱۳ هر جمله منفرد $p(x^{(t)} | x^{<t})$ احتمال شرطی نشانه $x^{(t)}$ به شرط وقوع (ظاهر شدن) t نشانه قبلی $x^{<t}$ در توالی است که به آن زمینه^۳ یا تاریخچه^۴ نیز می گویند. در عمل محاسبه این احتمال به صورت رابطه ۲-۱۳ تقریباً غیر ممکن است؛ زیرا، نیازمند داشتن همه توالی های ممکن هستیم. مدل های سنتی n-gram برای غلبه بر چالش های محاسباتی، با استفاده از فرض مارکوف رابطه ۲-۱۳ را به در نظر گرفتن تنها $n-1$ نشانه قبلی محدود می کنند. اگرچه در بسیاری از مسائل این مدل ها به خوبی پاسخگو هستند؛

¹Long-Short Term Memory

²Gated Recurrent Units

³Context

⁴History

اما، برای توالی‌های طولانی (بیشتر از ۴ یا ۵ نشانه) و نیز مشاهده نشده مناسب نیستند. در حالتی که نشانه فعلی پیش از این مشاهده نشده باشد، احتمال صفر به آن نسبت داده می‌شود که سبب صفر شدن احتمال پایانی می‌گردد. برای حل این مشکل، مدل‌های n-gram نیازمند اعمال فنون هموارسازی^۱ هستند [۳۹]. این فنون، راه‌حلی برای از بین بردن احتمال‌های صفر پیشنهاد می‌دهند. در هر صورت، مدل‌های n-gram با محدودیت‌های جدی روبه‌رو هستند.

۲-۴-۱ مدل زبانی عصبی

می‌توان از RNN برای ایجاد مدل زبانی استفاده کرد. این مدل‌ها را مدل زبانی عصبی (NLM) می‌گویند. استفاده از RNN هر دو مشکل اشاره شده در بالا را برطرف می‌کند. یعنی هم امکان پیش‌بینی توالی‌های طولانی‌تر فراهم می‌شود و هم وقوع احتمالات صفر از بین می‌رود [۴۵]. شکل ۲-۷ یک معماری از RNN برای ساخت مدل زبانی در NLP را نشان می‌دهد. چون هدف مدل زبانی پیش‌بینی نشانه بعدی است، توالی x را در ورودی شبکه قرار داده و خروجی متناظر با آن را همان توالی x که یک واحد روبه‌جلو شیفت داده شده است، می‌گذارند. این روش آموزش شبکه Teacher Forcing می‌گویند [۳۸].

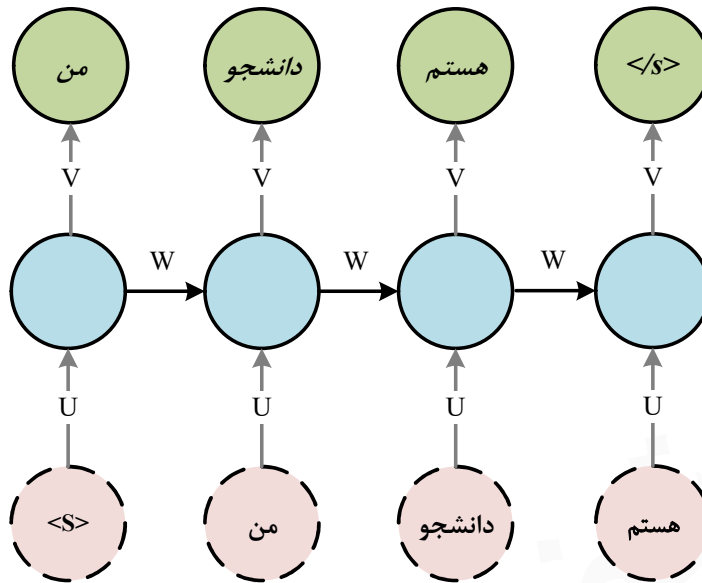
مدل‌های زبانی را مدل مولد نیز می‌نامند؛ زیرا یک توزیع احتمالی روی توالی‌های یک زبان می‌دهند که با نمونه‌برداری از آن می‌توان توالی‌های جدید تولید کرد. در NLP نشانه‌های هر توالی واژه‌ها هستند، اما می‌توان این مدل را در سطح کاراکتر نیز آموزش داد. ایده اصلی در این پایان‌نامه نیز همین است. هر فایل زبان و قواعد مخصوص به خود را دارد که در مشخصه‌های قالب فایل ذکر شده است. با ایجاد یک مدل زبانی برای این زبان در سطح کاراکتر، امکان پیش‌بینی نشانه بعدی از روی یک توالی از نشانه‌های داده شده فراهم می‌شود. با نمونه برداری از این مدل زبانی می‌توان داده‌های آزمون جدیدی ایجاد کرد.

۲-۴-۲ ارزیابی مدل زبانی

چگونه می‌توان میزان خوب بودن یک مدل زبانی را تعیین کرد؟ یک روش برای ارزیابی مدل زبانی، قرار دادن آن در یک وظیفه مشخص و اندازه‌گیری میزان دقت نتایج حاصله است. به این روش ارزیابی بیرونی^۲ می‌گویند [۳۹]. برای مثال در یک وظیفه تصحیح خودکار غلط‌های املایی، تعداد واژه‌های نادرستی که به‌درستی با واژه‌های صحیح جایگزین شده‌اند را شمارش می‌کنیم و بر تعداد کل واژه‌های غلط تقسیم می‌نماییم. ارزیابی

^۱Smoothing

^۲Extrinsic Evaluation



شکل ۲-۷: طرحواره‌ای از معماری یک مدل زبانی ایجاد شده با استفاده از RNN. $< s >$ نشانه شروع توالی و $< /s >$ نشانه خاتمه توالی در ابتدا و انتهای هر توالی موجود در مجموعه آموزش قرار داده می‌شود. بدین ترتیب، تولید توالی جدید پس از پیش‌بینی نشانه پایان، متوقف می‌گردد [۴۵].

بیرونی، زمان‌بر، پرهزینه و مختص به وظیفه اجرا شده خواهد بود. روش دیگر استفاده از ارزیابی درونی^۱ است [۳۹]. معیار سرگشتگی^۲ با هدف ارزیابی درونی مدل زبانی، به صورت زیر تعریف می‌شود [۴۶]:

$$\begin{aligned}
 PPLM(x) &= \sqrt[n]{\prod_{i=1}^n \left(\frac{1}{p(x^{(i)} | < x^{(1)}, \dots, x^{(i-1)} >)} \right)} \\
 &= 2^{-\frac{1}{n} \sum_{i=1}^n \log_2 p(x^{(i)} | < x^{(1)}, \dots, x^{(i-1)} >)}
 \end{aligned}
 \quad (۱۴-۲)$$

در رابطه ۲-۱۴، x توالی مورد ارزیابی است. همان‌طور که مشاهده می‌شود سرگشتگی در ارتباط مستقیم با خطای آنتروپی متقاطع (رابطه ۲-۶) است که میزان اختلاف بین مدل و نمونه‌های مجموعه آزمون را نشان می‌دهد. بنابراین هرچه قدر میزان سرگشتگی پایین‌تر باشد، مدل زبانی بهتر است. پایه توان و پایه لگاریتم در رابطه ۲-۱۴ می‌تواند مقادیری به غیر از ۲ انتخاب شود. اما در هر حالت هر دو پایه بایستی یکسان باشند تا تساوی اول (بالا) و دوم (پایین) رابطه برقرار گردد. معمولاً از پایه لگاریتم طبیعی در این رابطه استفاده

^۱Intrinsic Evaluation

^۲Perplexity

می‌شود [۴۷].

برای درک بهتر نحوه ارزیابی توسط این معیار، توالی $x = \langle x^{(1)}, x^{(2)}, x^{(3)}, \dots, x^{(n)} \rangle$ را در نظر می‌گیریم که متشکل از V نشانه متفاوت است. V مجموعه واژگان^۱ زبان است. در غیاب مدل زبانی (بدترین حالت)، تنها می‌توانیم ادعا کنیم که هر نشانه دست‌کم با احتمال $\frac{1}{V}$ در توالی رخ می‌دهد (ضریب انشعاب^۲). بدیهی است که این احتمال برای همه نشانه‌ها یکسان است. برای توالی x پس از جایگزینی احتمال آن در رابطه ۲-۱۴ داریم:

$$PP_{LM}(x) = \sqrt[n]{\prod_{i=1}^n \left(\frac{1}{V}\right)} = \sqrt[n]{V^n} = V$$

بنابراین در بدترین حالت میزان سرگشتگی برابر با اندازه مجموعه واژگان زبان است. حال چنانچه از مدل زبانی استفاده کنیم، احتمال نسبت داده شده به وقوع هر نشانه، نسبت به حالت عدم استفاده از مدل زبانی، افزایش و در نتیجه سرگشتگی کاهش خواهد یافت. در ارزیابی مدل‌های این پایان‌نامه از معیار سرگشتگی استفاده می‌کنیم.

۲-۵ خلاصه

در این فصل مفاهیم اولیه سه حوزه کلی مطرح در این پایان‌نامه یعنی آزمون نرم‌افزار، آزمون فازی و یادگیری ژرف را بیان کردیم. آزمون نرم‌افزار به عنوان یک مسئله تصمیم‌ناپذیر تعریف گردید و معیارهای پوشش برای ارزیابی آن معرفی شدند. آزمون فازی به عنوان یک فن مطرح در آزمون نرم‌افزار، معرفی و از سه دیدگاه مختلف مورد طبقه‌بندی و بحث قرار گرفت: نخست، از دیدگاه اطلاعات در دسترس از SUT به سه دسته آزمون جعبه سیاه، جعبه سفید و جعبه خاکستری تقسیم گردید. دوم، از دیدگاه روش‌های تولید داده آزمون به سه دسته مبتنی بر جابه‌جایی، مبتنی بر تولید و نیز ترکیبی تقسیم‌بندی شد. سوم، با توجه به دریافت بازخورد از نتیجه اجرا به دو دسته دارای بازخورد و بدون بازخورد تفکیک گشت. ابزار خودکارسازی آزمون فازی تحت عنوان کلی فازر معرفی و معماری آن تشریح شد. فازرها با توجه به طبقه‌بندی‌های فوق در بحث آزمون فازی و نیز با توجه به نوع ورودی SUT متفاوت هستند. فازرهای قالب فایل مربوط به برنامه‌هایی با ورودی فایل هستند. برای اطلاعات بیشتر در زمینه آزمون نرم‌افزار خواننده را به [۲۷، ۱۳] ارجاع می‌دهیم. در زمینه آزمون فازی نیز خواننده را به [۲۹، ۲۴، ۲] رجوع می‌دهیم.

در بخش پایانی این فصل، یادگیری ژرف به عنوان زیرشاخه‌ای از یادگیری ماشینی با استفاده از شبکه‌های

¹Vocabulary

²Branching Factor

عصبی مصنوعی، تفهیم شد. شبکه عصبی مورد استفاده در یادگیری ژرف را شبکه عصبی ژرف گویند. RNN نوع خاصی از شبکه‌های عصبی ژرف است که برای پردازش وظیفه‌های مبتنی بر توالی مثل مدل زبانی استفاده می‌شود. ساختار یک قالب فایل زبان مختص به خود را دارد و هدف از مطرح کردن اصول اولیه یادگیری ژرف استفاده از مدل‌های این فن، در یادگیری ساختار فایل، جهت تولید داده آزمون جدید است که در فصل ۴ به آن خواهیم پرداخت. الگوریتم‌های آموزش شبکه‌های عصبی ژرف ریاضیات گسترده‌ای دارد که به‌طور خلاصه به موارد مهم آن اشاره شد. جهت مطالعه مبسوط این الگوریتم‌ها خواننده را به [۳۸] ارجاع می‌دهیم. همچنین توضیح کاملی از روش‌های یادگیری توالی با استفاده از RNN‌ها در [۴۴] آمده است. در فصل بعد برخی از تازه‌ترین کارهای مرتبط با مفاهیم مطرح شده در این فصل توضیح داده می‌شوند.

فصل ۳

کارهای مرتبط

«یک شب تاریک و طوفانی بود. پاییز ۱۹۹۴، در آپارتمان خود در مادیسون نشسته بودم. آن شب من از طریق خط تلفن به سیستم‌های یونیکس دانشگاه متصل شدم. با بارش سنگین باران اختلال زیادی روی خط اتصالی بود که در اجرای فرمان‌های ارسالی من دخالت می‌کرد. رقابتی بین سرعت نوشتن یک فرمان قبل از آنکه اختلال آن را خراب کند وجود داشت. چیزی که مرا شگفت زده می‌کرد این حقیقت بود که اختلال سبب ایجاد خرابی و سقوط برنامه‌ها می‌شد و شگفت‌آورتر برنامه‌هایی بود که سقوط می‌کرد: ابزارهای رایج یونیکس!»

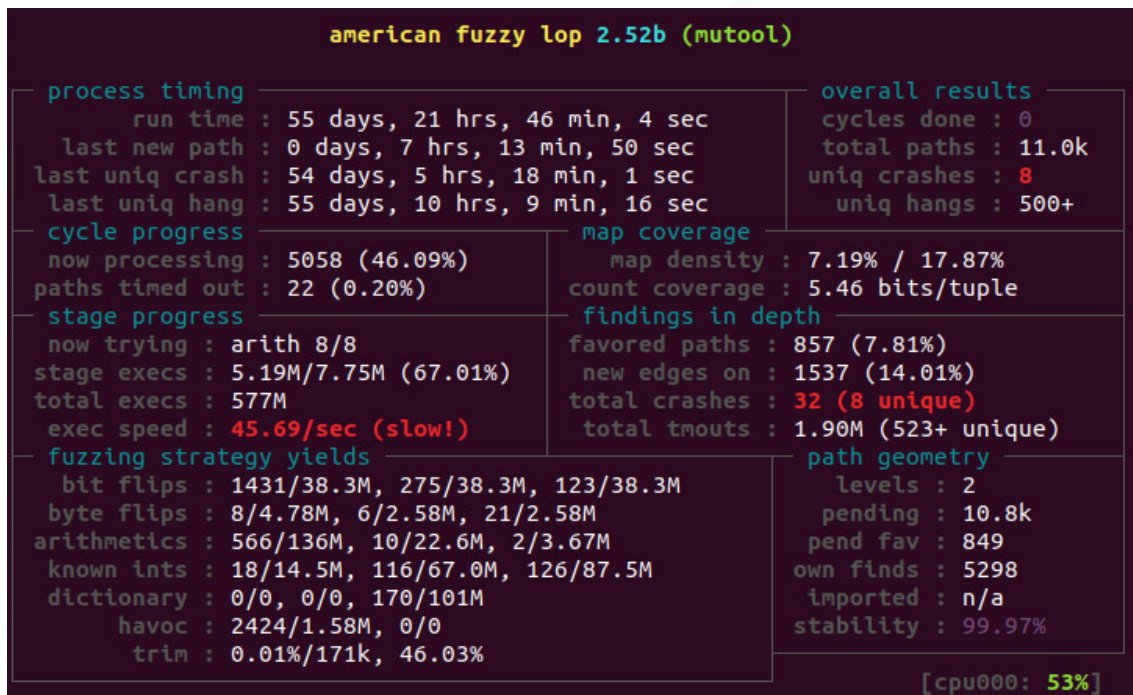
✠ بارتون میلر، مبدع آزمون فازی

کارهای بسیاری برای بهبود آزمون فازی اولیه که داده‌های آزمون را به صورت تصادفی تولید می‌کرد [۳]، انجام شده است. فازرهای مبتنی بر تولید در برنامه‌های با ساختار ورودی پیچیده پوشش کد بیشتری نسبت به فازرهای مبتنی بر جابه‌جایی فراهم می‌کنند [۳۴]، اما کاملاً خودکار نیستند [۱۱]. در مقابل فازرهای مبتنی بر جابه‌جایی سعی کرده‌اند تا با استفاده از الگوریتم‌های تکاملی مثل ژنتیک داده آزمون‌های بهتری را برای جابه‌جایی انتخاب کنند. AFL [۱۹] نمونه‌ای موفق از فازرهای مبتنی بر جابه‌جایی است. در سال ۲۰۱۷، فنون یادگیری ماشینی به هر دودسته از فازرهای بالا اعمال شده‌اند. در فازرهای مبتنی بر تولید، برای یادگیری خودکار گرامر فایل [۱۱] و در فازرهای مبتنی بر جابه‌جایی برای پیش‌بینی بهترین مکان جابه‌جایی [۲۰]. هریک از این کارها نواقص و محدودیت‌هایی دارند. در این فصل به معرفی، نقد و بررسی این کارها می‌پردازیم.

۳-۱ فازر AFL

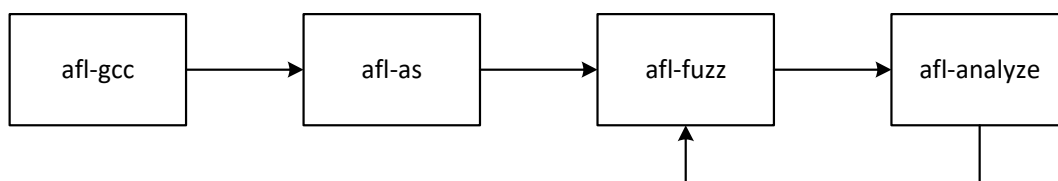
AFL [۱۹] یک فازر قالب فایل جعبه خاکستری، با تولید داده آزمون مبتنی بر جابه‌جایی، دارای بازخورد، متن‌باز و رایگان است که توسط Michal Zalewski^۱ توسعه داده شده است. این فازر روی سیستم عامل‌های خانواده یونیکس قابل اجرا است. راه‌اندازی آن ساده بوده و واسط کاربری خوبی برای دنبال کردن جزئیات آماری فرایند آزمون فازی و خطاهای کشف شده دارد. در شکل ۳-۱ یک نمونه از اجرای این فازر را در عمل نشان داده‌ایم.

به‌طور پیش‌فرض AFL برای سنجش پوشش کد، به کد منبع SUT نیاز دارد. برنامه‌های نوشته شده به زبان‌های C، C++ و Objective-C در آزمون فازی جعبه سفید با AFL قابل استفاده هستند. همچنین نسخه‌هایی از AFL برای برنامه‌های نوشته شده به زبان‌های Go و Python توسط دیگران انتشار یافته است. AFL در آزمون فازی جعبه سیاه، برای ابزارگذاری و اخذ اطلاعات زمان اجرا، از ابزار QEMU [۲۳] استفاده می‌کند. این فازر به‌حدی موفق بوده که پژوهش‌های زیادی برای بهبود جنبه‌های مختلف آن انجام شده است. اما چنانچه خواهیم دید روی قالب فایل‌هایی با ساختار پیچیده نمی‌تواند به پوشش خوبی دست پیدا کند [۲۰].



شکل ۳-۱: اجرای AFL بر روی ابزار mutool از نرم‌افزار MuPDF [۲۵]. این تصویر جزئیات اجرا بعد از گذشت ۵۵ روز از آغاز فرایند آزمون فازی را نشان می‌دهد. آزمون تا زمانی که کاربر `ctrl+z` را فشار ندهد، اجرا می‌شود.

^۱<http://lcamtuf.coredump.cx/>



شکل ۳-۲: مؤلفه‌های فازر AFL و ارتباط آنها با یکدیگر [۱۹].

۳-۱-۱ معماری AFL

شکل ۳-۲ مؤلفه‌های اصلی فازر AFL و ترتیب استفاده از آنها در هنگام آزمون فازی را نشان می‌دهد. در این معماری چهار مؤلفه مشاهده می‌شود:

- **afl-gcc**. یک جایگزین برای gcc یا clang استاندارد است که برای کامپایل کد منبع SUT استفاده می‌شود و خروجی آن به afl-as ارسال می‌شود. در رویکرد جعبه سفید استفاده مرحله اول کامپایل SUT با afl-gcc است.

- **afl-as**. کد کامپایل شده با afl-gcc را با تزریق کدهای اسمبلی ابزارگذاری می‌کند. ابزارگذاری به نحوی است که پوشش انشعاب یا پرش را ضبط می‌کند. خروجی afl-as فایل دودویی قابل اجرای SUT است که سپس توسط afl-fuzz استفاده می‌شود.

- **afl-fuzz**. همانطور که انتظار می‌رود، هسته اصلی فازر است که عملیات فاز ورودی را انجام می‌دهد. این ابزار فایل دودویی و داده آزمون ورودی را دریافت و با توجه به اطلاعات دریافتی از afl-analyze فرایند آزمون را ادامه می‌دهد. این مؤلفه برای آزمون فازی جعبه سیاه، مستقیماً به کار می‌رود. afl-fuzz همچنین مسئول چاپ واسط کاربری روی ترمینال است.

- **afl-analyze**. اثر ورودی اجرا شده بر پوشش کد را بررسی می‌کند و اطلاعات مربوط به پوشش کد را به عنوان بازخوردی به afl-fuzz می‌دهد تا برای جابه‌جایی‌های بهتر در ادامه استفاده کند.

حلقه بازخوردی که در شکل ۳-۲ وجود دارد نشان دهنده تکاملی بودن فرایند آزمون فازی در AFL است. AFL در هسته خود از الگوریتم ژنتیک استفاده می‌کند. یک فایل ورودی جهش یافته، مفید و مورد توجه است اگر قسمت‌های جدیدی از کد دودویی را اجرا کرده باشد یا تعداد اجرای کدهای قبلی مشاهده شده را افزایش داده باشد. این ویژگی با عنوان Input Gain شناخته می‌شود [۲۰]. AFL، سپس این داده آزمون را به انتهای

¹Component

صف داده‌های آزمون، اضافه می‌کند. به این ترتیب حاصل فرایند آزمون فازی در AFL، علاوه بر اجرای SUT و اندازه‌گیری پوشش کد، یک مجموعه داده آزمون جهش یافته با ویژگی‌های متمایز است.

برای مقایسه نحوه عملکرد AFL و یک فازر تصادفی، جزئیات الگوریتم‌های این دو فازر را مرور می‌کنیم. الگوریتم ۳-۱، فرایند فاز ورودی در یک فازر تصادفی و الگوریتم ۳-۲ همین فرایند را در AFL نشان می‌دهد. تابع *Mutate* یک بایت از ورودی را به‌صورت درجا، با استفاده از فنونی مثل وارون‌کردن بیت، وارون‌کردن بایت، چرخش بیت یا عملیات ریاضی و منطقی، جابه‌جا می‌کند. تابع *Execute* برنامه را با ورودی جابه‌جا شده اجرا و خرابی‌های احتمالی را گزارش می‌دهد. خطوط سایه زده شده در دو الگوریتم (خطوط ۱۰، ۱۴ و ۱۵) قسمت‌های متفاوت را نشان می‌دهد [۲۰].

الگوریتم ۳-۱ Basic-Random Fuzzing [۲۰]

Input: *Seeds*, Target program *P*

Output: *MaliciousInputs*

```

1 for Seed ∈ Seeds do
2   for iterations ← 0 to limit do
3     input ← Seed
4     length ← len(Seed)
5     mutations ← RandInt(length)
6     for mut ← 0 to mutations do
7       byte ← RandInt(length)
8       mutate(input, byte)
9     end
10    result ← Execute(P, input)
11    if result is crash then
12      Append input to MaliciousInputs
13    end
14
15
16  end
17 end
```


الگوریتم ۳-۲ AFL Fuzzing [۲۰]

Input: *Seeds*, Target program *P***Output:** *MaliciousInputs*

```

1 for Seed ∈ Seeds do
2   for iterations ← 0 to limit do
3     input ← Seed
4     length ← len(Seed)
5     mutations ← RandInt(length)
6     for mut ← 0 to mutations do
7       byte ← RandInt(length)
8       Mutate(input, byte)
9     end
10    result, cov ← Execute(P, input)
11    if result is crash then
12      Append input to MaliciousInputs
13    end
14    if HasInputGain(cov) then
15      Append input to Seeds
16    end
17  end
18 end

```

۳-۱-۲ مشکلات AFL

آزمون فازی از لحاظ محاسباتی سنگین است. یعنی حتی یک Input Gain کوچک نیاز نیازمند هزارها تا میلیون‌ها جابه‌جایی است [۲۰]. از طرفی همه جابه‌جایی‌های یکسان نیستند. AFL با بهره‌گیری از بازخورد داده‌های آزمون بهتری را انتخاب می‌کند، اما همچنان آنها را به صورت تصادفی جهش می‌دهد یا جابه‌جا می‌کند. در نتیجه تعداد زیادی داده آزمون تکراری تولید می‌شود که لزوماً تأثیری بر بهبود آزمون ندارند. از طرفی در ساختارهای پیچیده، تغییر برخی قسمت‌ها سبب می‌شود تا داده آزمون ورودی نامعتبر شود و سریعاً توسط تجزیه‌گر مردود اعلام گردد. در نتیجه تعداد زیادی داده آزمون هدر رفته خواهیم داشت. AFL^۱ افزوده [۲۰] مسئله انتخاب تصادفی مکان‌های جابه‌جایی را مورد مطالعه و راه‌کارهایی برای هوشمندسازی آن ارائه داده

^۱ Augmented

است (بخش ۲-۳).

محدودیت دیگر AFL قابلیت حمل^۱ پایین آن است. AFL به خاطر استفاده از فراخوانی های سیستمی سیستم عامل Linux تنها در توزیع های آن قابل استفاده است. نسخه ای از AFL، تحت عنوان WinAFL^۲ برای سیستم عامل ویندوز توسعه داده شده است که البته سرعت پایین و سربار بالایی دارد. WinAFL برای افزایش سرعت پیشنهاد می کند که تنها یک تابع از برنامه که هدف اصلی آزمون فازی است، ابزارگذاری شود. به دلیل مشکلاتی که اشاره شد، اجرای AFL و هرگونه فازر مشابه آن، روی نرم افزارهایی مثل PDF خوان ها که ساختار ورودی آنها پیچیده است به پوشش کد به مراتب کمتری دست می یابد که حتی با افزایش زمان آزمون نیز نتایج بهبود چندانی نمی کند. در این پایان نامه یک روش ترکیبی برای تولید داده آزمون ارائه می دهیم که پوشش کد بهتری نسبت به فازرهای مبتنی بر جابه جایی محض برای ساختارهای پیچیده فراهم کرده و در عین حال ویژگی های مثبت جابه جایی در فایل های دودویی را نیز به کار می بندد.

۲-۳ فازر AFL/افزوده

AFL/افزوده [۲۰] تلاش می کند با استفاده از فنون یادگیری ماشینی مکان های مناسبی را برای جابه جایی بایت ها پیدا کند. چارچوب ارائه شده در این کار شامل فازر AFL و یک مدل است که مکان های مفید برای جابه جایی را مشخص می کند. در طول اجرا فازر ابتدا مدل را برای اخذ آدرس مکان های جهش پرس و جو می کند و جابه جایی های فازر را روی مکان های بازگردانیده شده متمرکز می کند.

برای آموزش مدل فایل ورودی از مجموعه دانه اولیه، فایل جهش یافته و میزان پوشش کد هر دو فایل پس از اجرا توسط SUT مورد نیاز است. در فرایند آموزش اگر پوشش کد فایل جهش یافته و فایل والد یکسان باشد (یعنی میزان Input Gain صفر باشد)، مکان های جابه جا شده مکان های امیدبخشی نخواهند بود و چنانچه پوشش کد فایل جابه جا شده افزایش داشته باشد (یعنی میزان Input Gain بزرگتر از صفر باشد)، آنگاه جابه جایی ها امیدبخش محسوب می شوند. با تعریف یک تابع خطا که بسته به مقدار Input Gain امتیازهایی به هریک از دو حالت بالا نسبت می دهد، مدل در حالت بانظارت آموزش می بیند. ورودی مکان های جهش در فایل و خروجی میزان امتیاز کسب شده بر اثر هریک از این مکان ها است. صورت بندی ریاضی مطالب عنوان شده، به طور مفصل در [۲۰] ذکر شده است.

یک ویژگی مثبت این مقاله آموزش و استفاده از چندین مدل متنوع یادگیری ژرف و نیز آزمون چندین برنامه هدف در انجام آزمایش ها است که سبب می شود تأثیر مدل های مختلف را بتوان با یکدیگر مقایسه کرد. برای

^۱Portability

^۲<https://github.com/ivanfratric/winafl>

ایجاد مجموعه آموزش، فازر AFL با تعداد ۱۸۰ دانه اولیه برای هر SUT به مدت ۲۴ ساعت اجرا گردیده و اطلاعات لازم ضبط شده‌اند.

الگوریتم ۳-۳، فرایند فاز ورودی و انجام آزمون فازی در AFL/افزوده را نشان می‌دهد. خطوط سایه زده شده (خطوط ۲، ۱۱، ۱۲ و ۱۳) قسمت‌های اضافه شده به الگوریتم اصلی AFL هستند. چون جابه‌جایی‌های AFL در حالت عادی به صورت تصادفی است (خطوط ۷ و ۸ در الگوریتم ۳-۲ و خطوط ۸ و ۹ در الگوریتم ۳-۳)، تغییر انجام شده در AFL/افزوده بدین صورت است که ابتدا مکان‌های مناسب جابه‌جایی برای یک ورودی را توسط مدل پیش‌بینی می‌کند، سپس اجازه می‌دهد تا AFL ورودی را جابه‌جا کند. حال چنانچه مجموع شباهت جابه‌جایی‌های AFL با جابه‌جایی‌های گزارش شده که مطلوب مدل است، کمتر از یک حد آستانه α باشد (خط ۱۲ الگوریتم ۳-۳)؛ این داده آزمون تولیدی برای اجرا ارسال نمی‌گردد. محاسبه شباهت دو رشته بیتی نیز با استفاده از ترکیب عطفی آنها (خط ۱۲ الگوریتم ۳-۳) به راحتی امکان‌پذیر است. چنانچه دو رشته مشابه نباشند مجموع بیت‌های ترکیب عطفی آنها کمتر از مجموع بیت‌های هریک از آنها خواهد بود و بالعکس.

۳-۲-۱ مشکلات AFL/افزوده

الگوریتم AFL/افزوده روی چندین قالب فایل و چندین برنامه مختلف، مورد آزمایش قرار گرفته است؛ از جمله قالب‌های فایل PNG، XML و PDF. کمترین بهبود گزارش شده بازهم مربوط به قالب فایل PDF و نرم‌افزار MuPDF است. علت آن هم ساختار پیچیده و هم حجم بالای فایل‌های PDF است. میزان درصد پوشش کد گزارش شده برای نرم‌افزار MuPDF توسط AFL اصلی ۱۱/۶۳ و توسط AFL/افزوده در بهترین مدل ۱۱/۸۰ بوده که افزایش ناچیز ۰/۱۷ درصدی را نشان می‌دهد. در بخش نتایج این مقاله همچنین هیچ ارزیابی روی دقت مدل‌های مختلف آموزش داده شده انجام نشده و به یک نتیجه‌گیری کلی بسنده شده است.

یک تأکید مهم که در مستندات AFL به آن اشاره شده است کوچک نگاه داشتن فایل‌های مجموعه دانه اولیه است (توصیه شده است که از فایل‌هایی با حجم زیر ۱ کیلوبایت استفاده شود)؛ زیرا، در صورتی که از فایل‌های با حجم بالایی استفاده شود سرعت فازر به شدت کاهش می‌یابد. در واقع جابه‌جایی‌های بسیاری بایستی روی یک دانه اعمال گردد تا یک داده آزمون جدید ایجاد شود. این در حالی است که برخی ساختارها مثل PDF ذاتاً حجم بالایی دارند. مثلاً یک فایل PDF تنها حاوی عبارت Hello Wrold حجمی در حدود ۱ کیلوبایت دارد. برای این ساختارها، تولید یک فایل جدید از ابتدا، مثلاً از روی یک مدل، سرعت تولید داده آزمون را افزایش می‌دهد. استفاده از یک روش ترکیبی یعنی ثابت نگه داشتن بخش‌هایی از فایل و تولید بخش‌های دارای اهمیت در ساختار یک فایل سرعت تولید داده را حتی بیشتر از تولید یک فایل از ابتدا هم

الگوریتم ۳-۳ Augmented-AFL Fuzzing [۲۰]

Input: *Seeds*, Target program *P***Output:** *MaliciousInputs*

```

1 for Seed  $\in$  Seeds do
2   bytemask  $\leftarrow$  QueryModel(Seed)
3   for iterations  $\leftarrow$  0 to limit do
4     input  $\leftarrow$  Seed
5     length  $\leftarrow$  len(Seed)
6     mutations  $\leftarrow$  RandInt(length)
7     for mut  $\leftarrow$  0 to mutations do
8       byte  $\leftarrow$  RandInt(length)
9       Mutate(input, byte)
10    end
11    diff  $\leftarrow$  input  $\oplus$  Seed
12    if  $\sum \text{diff} \wedge \text{bytemask} < \alpha$  then
13      Continue
14    end
15    result, codecov  $\leftarrow$  Execute(P, input)
16    if result is crash then
17      Append input to MaliciousInputs
18    end
19    if HasInputGain(codecov) then
20      Append input to Seeds
21    end
22  end
23 end

```

افزایش می‌دهد. روش ارائه شده در این پایان‌نامه یک روش ترکیبی است.

۳-۳ یادگیری و فاز

مقاله یادگیری و فاز^۱ [۱۱] روشی برای تولید داده آزمون جهت استفاده در آزمون فازی بر مبنای RNN و مدل کدگذار-کدگشا^۲ [۱۷، ۱۸] ارائه کرده است. در این مقاله ساختار فایل PDF و مرورگر وب Edge^۳ شرکت مایکروسافت، برای آزمون انتخاب شده‌اند. ایده اصلی یادگیری یک مدل مولد زبان روی مجموعه‌ای از ویژگی‌های اشیای PDF با داشتن مجموعه‌ای از نمونه‌های اولیه است. ساختار قالب فایل PDF و اشیای داده‌ای آن، در پیوست آ توضیح داده شده است.

مدل کدگذار-کدگشا [۱۷، ۱۸] از دو RNN تشکیل شده است. یک شبکه کدگذار که یک توالی ورودی با بُعد متغیر را به یک نمایش بعد ثابت تبدیل می‌کند (بردار زمینه) و یک شبکه کدگشا که یک توالی ورودی بُعد ثابت را می‌گیرد و به یک توالی ورودی با بُعد متغیر تولید می‌کند. شبکه کدگشا توالی‌های خروجی را با استفاده از کاراکتر خروجی پیش‌بینی شده که در مرحله زمانی t به عنوان کاراکتر ورودی برای مرحله زمانی $t+1$ تولید شده است، تولید می‌کند. یک بازنمایی از معماری این مدل در شکل ۳-۳ نشان داده شده است. این مدل نخستین بار در وظیفه ترجمه ماشینی استفاده شده است [۱۸]. با توصیف بیان شده، این معماری امکان یادگیری توزیع شرطی $p(< y^{(1)}, y^{(2)}, y^{(3)}, \dots, y^{(n)} > | < x^{(1)}, x^{(2)}, x^{(3)}, \dots, x^{(n)} >)$ را فراهم می‌کند.

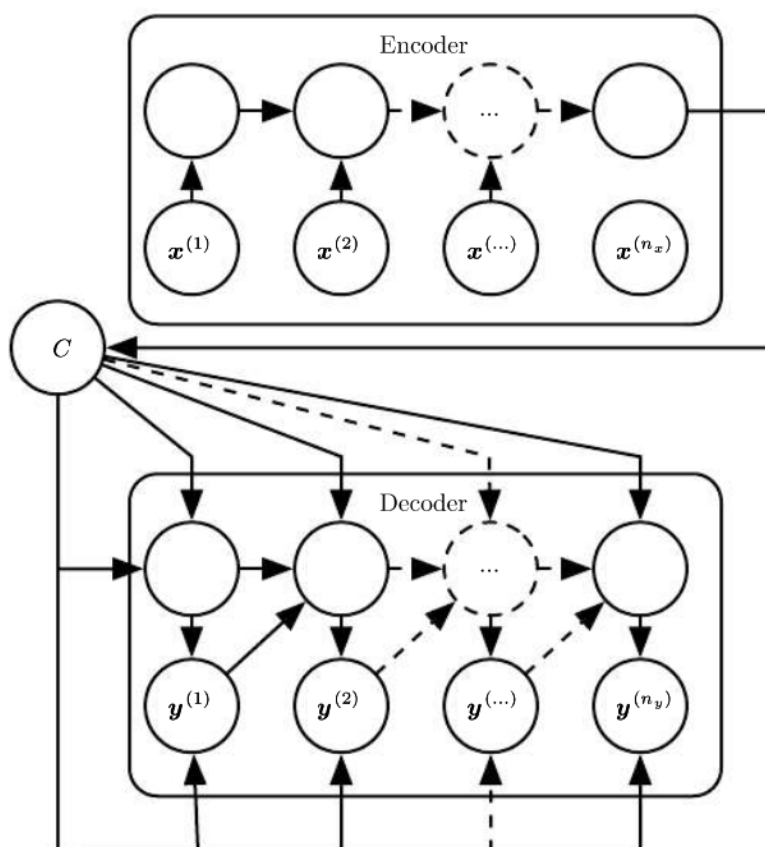
۳-۳-۱ آموزش مدل کدگذار-کدگشا

مدل کدگذار-کدگشا در روش یادگیری و فاز، با استفاده از تکه اشیای PDF که هر یک به صورت توالی‌ای از کاراکترها در نظر گرفته شده است، آموزش می‌بیند. برای آموزش، ابتدا همه فایل‌های حاوی اشیای PDF به شکل یک توالی از کاراکترها به یکدیگر الحاق می‌شوند که منجر به ایجاد یک توالی بزرگ از کاراکترها به صورت $\hat{s} = < s^{(1)}, s^{(2)}, s^{(3)}, \dots, s^{(n)} >$ خواهد شد. سپس توالی $\hat{s}$ به چند زیر توالی آموزشی با طول ثابت d شکسته می‌شود به نحوی که i امین نمونه آموزشی عبارت است از: $t_i = \hat{s}[i * d : (i + 1) * d]$ که $s[l : u]$ یک زیرتوالی از s بین اندیس‌های l و u را نشان می‌دهد. در ادامه، توالی خروجی برای هر توالی آموزشی t_i عبارت است از توالی ورودی که یک واحد به سمت راست شیفت داده شده است؛ یعنی:

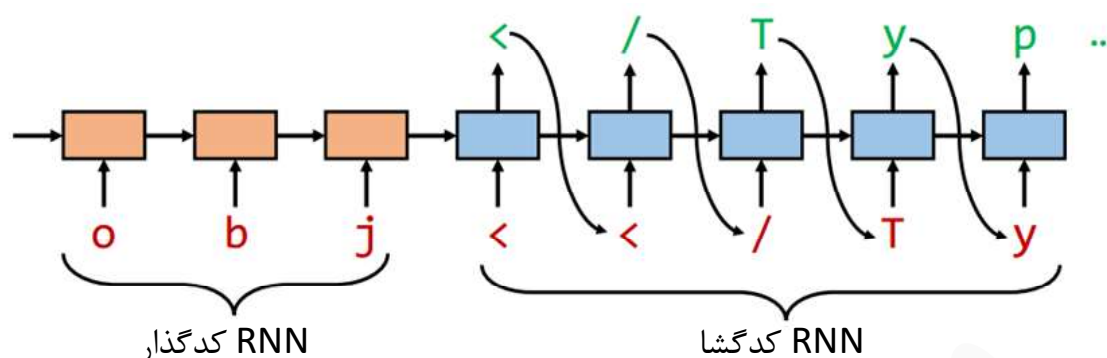
¹Learn&Fuzz

²Encoder-Decoder

^۳مرورگر جدید شرکت مایکروسافت که همراه با سیستم عامل ویندوز ۱۰ معرفی شد.



شکل ۳-۳: مثالی از معماری مدل کدگذار-کدگشا، متشکل از دو RNN و بردار زمینه C . این مدل برای یادگیری تولید توالی خروجی $\langle y^{(1)}, y^{(2)}, y^{(3)}, \dots, y^{(n_y)} \rangle$ از روی بردار زمینه، که نماینده توالی ورودی $\langle x^{(1)}, x^{(2)}, x^{(3)}, \dots, x^{(n_x)} \rangle$ است، به کار می‌رود [۳۸].



شکل ۳-۴: مدل RNN کدگذار-کدگشا برای تولید اشیای داده‌ای PDF [۱۱].

از آنجایی که مدل بحث شده در حالت بدون نظارت آموزش می‌بیند تا یک مدل مولد روی مجموعه همه توالی‌های ایجاد شده، یادگیری شود. شکل ۳-۴ مدل ارائه شده در این مقاله را نشان می‌دهد. $o_{t_i} = \hat{s}[i * d + 1 : (i + 1) * d + 1]$. مدل سپس به صورت انتها به انتها^۱ آموزش می‌بیند تا یک مدل مولد روی از آنجایی که مدل بحث شده در حالت بدون نظارت آموزش می‌بیند، برچسب‌های تولید شده به طور صریح برای مشخص کردن این که مدل یادگرفته شده چه قدر خوب عمل می‌کند، آزموده نشدند. به جای آن چندین مدل آموزش داده می‌شود که در تعداد دوره‌ها (بخش ۲-۳-۳) متفاوت هستند. مدل در ۱۰، ۲۰، ۳۰، ۴۰ و ۵۰ دوره آموزش داده شده است. در تنظیمات شبکه ارائه شده در این مقاله، یادگیری روی هر دوره حدود ۱۲ دقیقه طول می‌کشد و بنابراین مدل با ۵۰ عدد حدوداً ۱۰ ساعت زمان نیاز دارد تا کامل شود. شبکه به صورت RNN با ۲ لایه مخفی که هر لایه ۱۲۸ سلول LSTM دارد، است.

۳-۳-۲ تولید داده‌های جدید

از مدل یادگیری شده برای تولید اشیای جدید PDF استفاده شده است. راهبردهای مختلفی برای تولید اشیا، بسته به راهبرد نمونه‌برداری که برای نمونه‌برداری از توزیع یادگیری شده به کار می‌رود، وجود دارند. در این مقاله، با یک پیشوند^۲ از توالی "obj" (که آغاز یک شی را مشخص می‌کند) شروع و سپس مدل را پرس‌وجو کرده تا یک توالی از کاراکترهای خروجی تولید گردد، تا زمانی که مدل "endobj" را که مشخص کننده انتهای یک شی است، تولید نماید. سپس این مقاله، سه راهبرد مختلف را برای تولید اشیای جدید از روی مدل، استفاده کرده است [۱۱]:

^۱End to End

^۲Prefix

• **NoSample**. در این راهبرد تولید، کاراکتر پیش‌بینی شده با بالاترین احتمال به‌عنوان کاراکتر بعدی انتخاب می‌شود که در واقع انتخابی حریصانه^۱ است. این راهبرد برای تولید اشیای PDF ای که خوش‌شکل^۲ و سازگار^۳ هستند، کاملاً نتیجه‌بخش است ولی تعداد اشیائی را که می‌توان تولید کرد، محدود می‌کند. با دادن یک پیشوند مثل "obj" بهترین توالی از کاراکترهای بعدی به‌صورت یکتا مشخص می‌شود و بنابراین این راهبرد همان شی را نتیجه می‌دهد. این محدودیت مغایر با اهداف آزمون فازی بوده و مانع مفید بودن استفاده از آن در آزمون فازی می‌شود.

• **Sample**. در این راهبرد تولید، به‌جای انتخاب محتمل‌ترین کاراکتر پیش‌بینی شده، از توزیع یادگرفته شده، برای نمونه‌برداری کاراکتر بعدی در توالی‌ای که پیشوند آن داده شده است، استفاده می‌شود. راهبرد نمونه‌برداری قادر به تولید مجموعه گوناگونی از اشیا با ترکیب الگوهای مختلفی که از اشیای موجود یادگرفته شده است، خواهد بود. به دلیل نمونه‌برداری اشیای تولید شده، تضمینی نیست که آنها به میزان راهبرد قبلی، خوش‌شکل باشند، که این ویژگی از منظر آزمون فازی مفید است.

• **SampleSpace**. این راهبرد ترکیب دو راهبرد قبلی است. SampleSpace توزیع احتمالی برای تولید کاراکتر بعدی را تنها زمانی که پیشوند توالی کنونی، با فضای خالی خاتمه می‌یابد، نمونه‌برداری می‌کند. برای میان توکن^۴ ها (یعنی پیشوندهایی که با فضای خالی خاتمه نمی‌یابند)، بهترین کاراکتر را انتخاب می‌کند (راهبرد NoSample) است که در مقایسه با راهبرد Sample، ورودی‌های خوش‌شکل‌تری تولید می‌شود. چون نمونه‌برداری در پایان کاراکترهای فضای خالی صورت می‌گیرد؛ نمونه‌های جدید نیز در این راهبرد، ساخته خواهند شد.

۳-۳-۳ اعمال آزمون فازی

برای فاز داده‌های آزمون تولید شده، مقاله الگوریتم جدیدی تحت عنوان SampleFuzz (الگوریتم ۳-۴)، ارائه کرده است. SampleFuzz توزیع یادگیری شده $D(x, \theta)$ ، احتمال فازینگ یک کاراکتر t_fuzz و احتمال آستانه p_t که تصمیم می‌گیرد آیا کاراکتر پیش‌بینی شده اصلاح شود یا نه، را به عنوان ورودی می‌پذیرد. سپس برای توالی خروجی seq ، الگوریتم $D(x, \theta)$ را نمونه‌برداری می‌کند تا کاراکتر بعدی یعنی c و احتمال آن $p(c)$ را در یک مرحله زمانی مشخص t به‌دست آورد. اگر احتمال $p(c)$ از آستانه فراهم شده توسط کاربر (p_t)

¹Greedy

²Well-formed

³Consistent

⁴Token

بالا تر باشد؛ یعنی، اگر مدل مطمئن باشد که کاراکتر بعدی به احتمال زیاد c است، الگوریتم تصمیم می‌گیرد تا یک کاراکتر دیگر c' که احتمال کمینه $p(c')$ را دارد، جایگزین c کند. البته این اصلاح (فاز) تنها در صورتی که عدد تصادفی تولید شده p_fuzz از ورودی t_fuzz مقدار بیشتری داشته باشد، انجام می‌شود، که بدین ترتیب اجازه می‌دهد تا کاربر بتواند بر روی احتمال (درصد) کاراکترهای فاز شده کنترل داشته باشد.

الگوریتم ۳-۴ SampleFuzz [۱۱]

Input: $D(x, \theta), t_fuzz, p_t$ **Output:** seq

```

1  $seq \leftarrow "obj"$ 
2 while  $\sim seq.EndsWith("endobj")$  do
3    $c, p(c) \leftarrow sample(D(seq, \theta))$  /* Sample  $c$  from the learnt distribution */
4    $p\_fuzz \leftarrow Random(0, 1)$  /* Random variable to decide whether to fuzz */
5   if  $p\_fuzz > t\_fuzz \wedge p(c) > p\_t$  then
6      $c \leftarrow argmin_c \{p(c') \sim D(seq, \theta)\}$  /* Replace  $c$  by  $c'$  (with lowest likelihood) */
7   end
8    $seq \leftarrow seq + c$ 
9   if  $len(seq) > MaxLen$  then
10     $seq \leftarrow "obj"$  /* Reset the sequence */
11  end
12 end
13 Return  $seq$ 

```

الگوریتم SampleFuzz اطمینان می‌یابد که طول شی با عدد $MaxLen$ محدود شود؛ اما، شرط حلقه $while$ الگوریتم تضمین نمی‌کند که همواره خاتمه یابد. با این حال، نویسندگان گزارش کرده‌اند که در آزمایش‌های انجام شده در عمل، الگوریتم همواره پایان یافته‌است. چیدمان آزمایش‌ها و نیز تحلیل جامع نتایج در [۱۱] آمده است که نشان از برتری الگوریتم SampleFuzz در مقایسه با روش تصادفی و افزایش نسبی پوشش کد در سطح دستورات برنامه است.

۳-۳-۴. مشکلات روش یادگیری و فاز

مقاله یادگیری و فاز برای نخستین بار از یادگیری ماشینی در تولید داده آزمون فازی استفاده کرده است. در عین حال روش ارائه شده در این مقاله دارای مشکلات و کاستی‌هایی است که عبارتند از:

۱. همواره با یک پیشوند ثابت شروع به تولید داده جدید می‌کند که این سبب می‌شود تنوع کمتری حاصل شود. در نتیجه ناچار به پیچیده کردن فرایند نمونه برداری از مدل می‌شود. همچنین این پیشوند کوتاه است و حاوی اطلاعات بعدی یک شی داده‌ای نیست. می‌توان مدل‌هایی با قابلیت پذیرفتن پیشوندهای متفاوت ایجاد کرد.

۲. مدل کدگذار-کدگشا معمولاً برای نگاشت توالی‌های با طول و دامنه‌های متفاوت به یکدیگر، استفاده می‌شود. بارزترین مثال وظیفه ترجمه ماشینی است که هدف آن نگاشت جمله‌های یک زبان به زبان دیگر است [۱۷]. در این وظیفه، طول جمله زبان مقصد لزوماً برابر طول معادل همان جمله در زبان مبدأ نیست. به همین خاطر نیازمند معماری پیچیده کدگذار-کدگشا هستیم. وظیفه یادگیری خودکار ساختار فایل بیشتر به ایجاد یک مدل زبانی روی الفبای زبان فایل ورودی شباهت دارد که در نتیجه استفاده از یک مدل زبانی عصبی در این وظیفه، منطقی‌تر و احتمالاً نتیجه‌بخش‌تر به نظر می‌رسد. چنانچه در روش یادگیری و فاز هم می‌بینیم که ساختار و وظایف شبکه کدگذار و کدگشا به خوبی تفکیک و تفهیم نگردیده است و صرفاً به پیچیده‌تر شدن مدل انجامیده است.

۳. فقط داده‌های متنی فایل تولید و فاز شده‌اند در حالی که یک فایل باینری در حالت کلی حاوی ترکیبی از داده‌های متنی و غیرمتنی است. می‌توان یک روش ترکیبی برای این منظور ارائه داد.

۴. الگوریتم فاز ارائه شده، همان‌طور که دیدیم، ممکن است هیچ‌گاه پایان نیابد یا اجرای الگوریتم برای تولید یک نمونه خیلی طولانی شود که خوب به نظر نمی‌رسد. می‌توان با تعریف شرایطی از این اتفاق جلوگیری کرد. خاتمه‌یافتن در هر صورت، یک ویژگی است که در تعریف اصطلاح الگوریتم وجود دارد و از این بابت این مسئله بسیار مهم تلقی می‌شود.

۵. معماری مدل پیشنهادی به نحوی است که توالی‌های تولید شده برای آموزش در برگیرنده همه اطلاعات و وابستگی‌های ساختار فایل نیستند. در هنگام تولید توالی‌های آموزشی هر بار به اندازه d واحد از روی مجموعه داده پرس می‌کنیم. اگر d کوچک باشد، توالی‌های کوچکی خواهیم داشت که لزوماً شامل یک وابستگی طولانی نیستند و اگر d بزرگ انتخاب شود، تعداد وابستگی‌های کمتری تشخیص داده می‌شود. به این نکته بایستی توجه کرد که وابستگی‌های فایل‌های با قالب مختلف، یکسان نیستند. می‌توان راهکارهای بهتری در نحوه تولید ورودی و خروجی شبکه یافت.

۶. در هر بار اجرا فقط یک نمونه را تولید می‌کند و بنابراین نمونه‌های تولید شده در چندین اجرای متوالی کاملاً مستقل از یکدیگر هستند. ممکن است چند نمونه متوالی به هم وابستگی داشته باشند در این صورت بایستی بتوان به نحوی این پارامتر را در تولید نمونه‌ها گنجانند.

۷. و بالاخره در مقاله یادگیری و فاز تنها از یک مدل یادگیری ژرف برای تولید داده آزمون استفاده شده است که می‌توان با انتخاب مدل‌های مختلف یا تغییر در ابرپارمترهای در مورد تأثیر نوع مدل نیز بحث کرد و آزمایش‌های جدید را تجربه نمود.

۳-۴ دیگر کارهای مرتبط

تعداد دیگری کار در زمینه تولید داده آزمون در فازرها و به‌ویژه در زمینه یادگیری ساختار فایل انجام شده است. در فصل اول، به‌طور خلاصه به برخی از آنها اشاره کردیم [۱۵، ۱۶]. بیشتر این روش‌ها بسیار پیچیده بوده، درک و پیاده‌سازی آنها مشکل است و سربار زیادی را به عملیات آزمون فازی یا عملیات پیش‌پردازش، عملیات پیش از انجام آزمون فازی، تحمیل می‌کنند. در ضمن اغلب آنها در نهایت محدود به ساختارهای ورودی خاصی می‌شوند. تمرکز ما بر روی روش‌هایی است که کاملاً عملیاتی بوده و قابل استفاده بر روی نرم‌افزارهای دنیای واقعی هستند.

در آزمایشگاه تحقیقاتی مهندسی معکوس دانشگاه علم و صنعت ایران، تعدادی زیادی پروژه، کار پژوهشی و پایان‌نامه در ارتباط با آزمون فازی و تولید داده آزمون، نگارش شده است^۱. یعقوبی [۴۸] یک روش خودکار تولید داده آزمون با استفاده از الگوریتم ژنتیک برای آزمون فازی جعبه سیاه مرورگرهای وب، ارائه داده است. هدف الگوریتم در این فرایند تکاملی، تشخیص مجموعه‌ای از برچسب‌های HTML است که حداکثر درهم ریختگی حافظه و کاهش سرعت اجرایی مرورگر را موجب شوند. در واقع دسته برچسب‌های HTML به‌عنوان صفحات ورودی هر مرورگر، طوری ساخته می‌شوند که در هر مرحله نسبت به دسته برچسب‌های قبل، میزان مصرف حافظه اصلی را بالا ببرند.

امکان بهره‌گیری از این روش برای آزمون فازی و شناسایی خطاها و آسیب‌پذیری‌های دیگر نرم‌افزارهایی که کد منبع آنها در دسترس نیست نیز وجود دارد. اما تمرکز این روش روی تولید برچسب‌های HTML بوده و سرعت همگرایی آن کند است. علاوه بر این، تنها هدف الگوریتم ژنتیک در اینجا، افزایش حافظه تعریف

^۱ فهرست کاملی از کارهای مرتبط انجام شده، از طریق وب‌سایت رسمی آزمایشگاه به نشانی <http://parsa.iust.ac.ir> قابل دسترسی هستند. همچنین فهرستی سازمان‌دهی شده و کامل در ارتباط با منابع مختلف آزمون فازی در نشانی <https://github.com/secfigo/Awesome-Fuzzing> وجود دارد.

شده است در حالی که می‌دانیم صرف افزایش حافظه منجر به کشف خطا نمی‌شود. در سیستمی با حافظه پایین ممکن است سیستم عامل برنامه را پیش از سقوط متوقف سازد یا برنامه فقط بر اثر کمبود حافظه نتواند یک ورودی به اندازه کافی بزرگ را پردازش کند. در ساختارهای پیچیده که تجزیه‌گر سخت‌گیری دارند، تولید داده به این روش بسیار زمان‌بر است؛ زیرا، در عمل خیلی از داده‌های آزمون تولیدی مورد پردازش قرار نمی‌گیرند و حافظه‌ای به آنها تخصیص داده نمی‌شود.

امینی [۴۹] یک روش تکاملی با هدف حل مشکل بزرگی فضای ورودی، برای تولید داده آزمون ارائه داده است. در روش‌هایی مشابه این روش مسئله تولید داده آزمون به صورت یک مسئله جست‌وجو تعریف می‌گردد و سپس از الگوریتم‌های اکتشافی برای یافتن یک حل بهینه استفاده می‌شود. در راهکار امینی، زمانی که یک داده آزمون یک الگوی آسیب‌پذیر^۱ را شناسایی می‌کند، داده آزمون بعدی به گونه‌ای تولید می‌شود که منجر به افشای آن آسیب‌پذیری گردد. در نتیجه احتمال کشف آسیب‌پذیری‌های شناسایی شده بهبود خواهد داشت. الگوی آسیب‌پذیر در سطح کد منبع برنامه قابل شناسایی بوده و بنابراین این روش در فازهای جعبه سفید قابل استفاده است. البته می‌توان آن را به فازهای جعبه خاکستری هم تعمیم داد. روش مذکور روی برنامه‌های دنیای واقعی ارزیابی نشده است. همچنین یافته‌های آزمایش‌های آن، با هیچ روش هوشمند تولید داده آزمونی، مقایسه نگردیده است.

۵-۳ خلاصه

جدول ۳-۱، یک جمع‌بندی از مزایا و معایب سه کار اصلی معرفی شده در این فصل [۱۱، ۱۹، ۲۰] را در کنار یکدیگر نشان می‌دهد. همچنین فازر FileFuzz [۲]، که یک فازر قالب فایل تصادفی و بسیار اولیه است به عنوان یک فازر مبنایی برای مقایسه، در جدول آورده شده است. البته تا به امروز، فازرهای بسیار زیادی در زمینه آزمون فازی قالب فایل توسعه داده شده‌اند که مجال بررسی همه آنها در این پایان‌نامه نبوده و در اینجا صرفاً آن دسته از کارهایی را بیان کردیم که قصد داریم در روش پیشنهادی خود مشکلات موجود در آنها را تا حد امکان برطرف کنیم. برخی از این فازرها مانند [۱۰] به صورت داخل سازمانی و یا تجاری ارائه شده‌اند و کد منبع آنها در دسترس نیست. عدم دسترسی به کد منبع یک فازر و در نتیجه عدم امکان تغییر در پیمانه‌های مختلف آن را می‌توان انگیزه‌ای برای توسعه فازرهای جدید دانست.

به‌طور خلاصه، در این فصل ابتدا فازر قالب فایل AFL [۱۹] را به عنوان یک فازر با روش تولید داده آزمون مبتنی بر جابه‌جایی مورد مطالعه قرار دادیم و دیدیم که به علت جابه‌جایی تصادفی بایت‌های فایل، تعداد زیادی داده آزمون بی‌هوده برای ساختارهای پیچیده ایجاد می‌کند. سپس AFL افزوده [۲۰] را توضیح

^۱ Vulnerable Pattern

جدول ۳-۱: مقایسه کارهای مرتبط

فازر	مشخصه‌ها	مزایا	معایب
FileFuzz [۲]	• مبتنی بر جابه‌جایی • جعبه سیاه	• سادگی پیاده‌سازی • عدم نیاز به اطلاع داشتن از ساختار ورودی و ساختار SUT	• تعیین مکان جابه‌جایی‌ها به صورت کاملاً تصادفی • پوشش کد بسیار پایین و غیر قابل قبول، برای بیشتر نرم‌افزارها پس از آزمون
AFL [۱۹]	• مبتنی بر جابه‌جایی • مبتنی بر بازخورد • جعبه خاکستری	• بیشینه کردن پوشش کد با بهره‌گیری از الگوریتم ژنتیک • قابل استفاده به صورت جعبه سیاه، جعبه خاکستری و جعبه سفید	• تعیین مکان جابه‌جایی‌ها به صورت تصادفی • تولید تعداد زیادی داده آزمون هدر رفته • سرعت پایین فاز در فایل‌های با حجم بیشتر از ۱ کیلوبایت
AFL افزوده [۲۰]	• مبتنی بر جابه‌جایی • مبتنی بر بازخورد • جعبه خاکستری	• تعیین مکان (آدرس نسبی) جابه‌جایی‌ها از طریق یک مدل ارزش‌دهی به بایت‌ها	• پوشش کد پایین برای نرم‌افزارهای با ساختار ورودی پیچیده
Learn&Fuzz [۱۱]	• مبتنی بر تولید • جعبه سفید	• یادگیری گرامر فایل و تولید داده‌های آزمون جدید با استفاده از مدل کدگذار-کدگشا • تعیین مکان فاز داده آزمون توسط مدل یادگیری شده • تعیین مقدار فاز شده داده آزمون	• حذف کامل داده‌های دودویی • شروع تولید داده آزمون با یک پیشوند ثابت • عدم تضمین پایان یافتن الگوریتم فاز • در دسترس نبودن فازر، مجموعه داده و دیگر پارامترهای مهم مربوط به آموزش مدل

دادیم از که فنون یادگیری ماشینی برای تعیین محل مکان‌های جابه‌جایی استفاده می‌کند. ایده اصلی در این روش، آن است که در قالب‌های فایل، معمولاً فقط قسمت‌های کوچکی از فایل برای تجزیه‌گر مهم هستند و در نتیجه بیشتر جابه‌جایی آن‌ها منتهی به اجرای مسیرهای جدید SUT می‌گردد. این روش نیز روی قالب پیچیده‌ای مانند PDF بهبود ناچیزی داشته است. با این حال روی قالب‌های کاملاً دودویی مثل قالب فایل PNG در زمان برابر به پوشش کد بهتری دست یافته است.

در ادامه فصل، روش یادگیری و فاز [۱۱] را توضیح دادیم که از مدل کدگذار-کدگشا برای تولید و فاز داده‌آزمون استفاده می‌کند. ایده اصلی این روش یادگیری خودکار گرامر برای یک قالب فایل متنی و سپس تولید داده با روش مبتنی بر تولید است. این روش با یک روش تصادفی که در ابتدای فصل نیز مطرح شد، مقایسه شده و نتایج بهبود داشته است. روش یادگیری و فاز روی ابزار متن بسته مرورگر Edge شرکت مایکروسافت و توسط محققین همین شرکت مورد آزمایش قرار گرفته و هیچ‌گونه کد منبع و مجموعه داده‌ای از آن منتشر نشده است. در مقابل اما فازر AFL متن‌باز و رایگان است. AFL افزوده نیز به‌تازگی توسط توسعه‌دهندگان آن منتشر شده است.

در پایان این فصل مروری اجمالی بر دو کار پیشین انجام شده در زمینه آزمون فازی و تولید خودکار داده آزمون در آزمایشگاه مهندسی معکوس دانشگاه علم و صنعت ایران، که این پایان‌نامه در آنجا نگارش یافته است، داشتیم. هر دو کار [۴۸، ۴۹] از الگوریتم ژنتیک برای تولید داده‌های آزمون استفاده کرده بودند که مزایا و معایب هریک نیز بیان شد. در این پایان‌نامه ما از الگوریتم‌های یادگیری ژرف برای تولید خودکار داده آزمون استفاده خواهیم کرد. با توجه به مزایا و معایبی که در کلیه روش‌های مطرح در این فصل شناسایی و معرفی کردیم، در فصل بعد یک روش ترکیبی که بتواند از ویژگی‌های خوب جابه‌جایی تصادفی فایل‌های دودویی در کنار تولید فایل‌های متنی براساس مدل، استفاده کند، ارائه خواهیم کرد. هدف بهبود برخی کاستی‌های شناسایی شده در کارهای مطرح در این فصل و حل مسائل اصلی اشاره شده در فصل اول است.

فصل ۴

روش پیشنهادی

«اولین قانون هر فناوری مورد استفاده در یک کسب و کار این است که اعمال خودکارسازی به یک عملیات کارآمد، کارآمدی را افزایش می‌دهد. دومین قانون این است که اعمال خودکارسازی به یک عملیات ناکارآمد، ناکارآمدی را افزایش می‌دهد.»

❖ بیل گیتس

هدف، در روش پیشنهادی همان‌طور که پیش از این هم گفته شد، ایجاد یک مدل برای یادگیری خودکار گرامر قالب فایل و سپس استفاده از آن در تولید داده‌های آزمون جدید است. معماری مدل یادگیرنده ساختار فایل، نحوه آموزش، مجموعه داده لازم برای آموزش مدل، نحوه تولید داده‌های آزمون و در نهایت چگونگی انجام آزمون فازی، مباحثی هستند که در این فصل به آنها خواهیم پرداخت. در ابتدا به بیان کلیات روش پیشنهادی پرداخته و سپس هریک از مراحل را با جزئیات تشریح می‌کنیم.

۴-۱ کلیات

شکل ۴-۱ یک روندنمای ساده از کلیات مراحل انجام روش پیشنهادی ما را نشان می‌دهد. در ساختار برخی قالب‌های فایل پیچیده مانند PDF، بسیاری از قسمت‌ها به صورت متنی هستند. اما ممکن است قسمت‌های دودویی هم وجود داشته باشد. در روش پیشنهادی خود ما ابتدا کلیه قسمت‌های متنی فایل‌های موجود در

دانه‌های اولیه را استخراج، سپس این قسمت‌ها را به یکدیگر الحاق می‌کنیم. حاصل این کار ایجاد یک پیکره^۱ بزرگ است که در بطن خود مشخصه‌های قالب یک فایل را دارد. در ادامه این مجموعه را مطابق با ضوابط فنون یادگیری ماشینی به سه مجموعه آموزش، ارزیابی و آزمون تقسیم‌بندی با نسبت‌های مشخص می‌کنیم. سپس یک مدل مولد را ایجاد و آن را روی داده‌های مجموعه آموزش، با تعریف یک روش یادگیری، آموزش می‌دهیم. در نهایت از این مدل برای ایجاد داده‌های آزمون جدید و اعمال آزمون فازی استفاده خواهیم کرد. اما چون مدل صرفاً قادر به تولید داده‌های متنی است در هنگام ایجاد یک داده آزمون، با فونونی که در ادامه خواهیم گفت، قسمت‌های دودویی را نیز به داده آزمون تزریق می‌نماییم.

۴-۱-۱ تمایز داده‌های متنی و دودویی

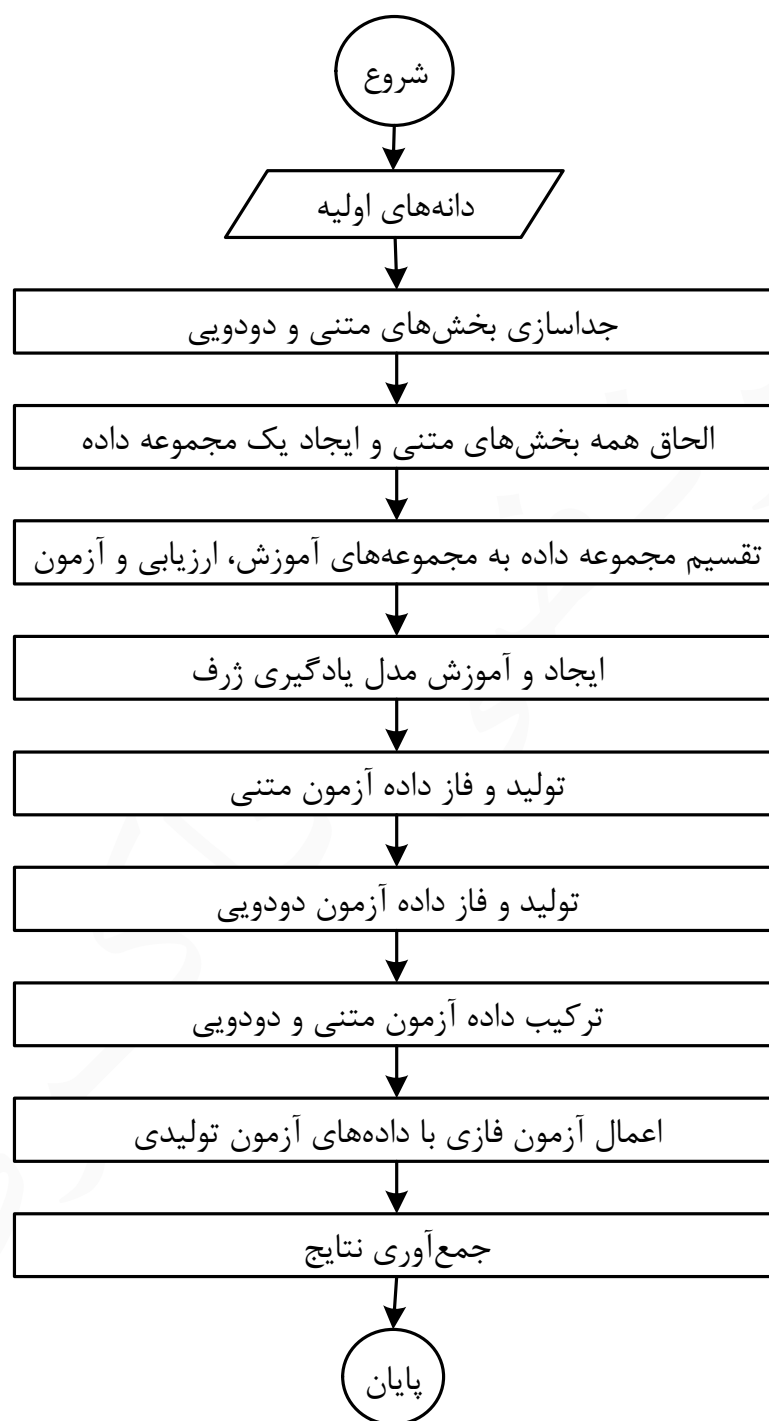
گفتیم مدل مولد را روی بخش متنی قالب فایل ایجاد می‌کنیم. اما مفهوم دقیق متنی و دودویی و معیار جداسازی آنها چیست و دلیل تأکید ما بر این موضوع از کجا نشئت می‌گیرد؟ منظور از داده‌های متنی در اینجا مجموعه کدهای ASCII^۲ است. هنگامی که در سطح بیت یک فایل را بررسی می‌کنیم تمایزی میان واژه‌های متنی و دودویی نیست؛ زیرا همه چیز با مفهوم صفر و یک بیان می‌شود که دودویی است. اما اگر فایل را به صورت یک توالی از بایت‌ها در نظر بگیریم، با توجه به اینکه هر بایت ۲۵۵ حالت ممکن خواهد داشت، می‌توانیم آن فایل را یک زبان متشکل از تکرار حداکثر ۲۵۵ نشانه بدانیم. کدهای ASCII حتی این تعداد را به ۱۲۸ نشانه که کاراکترها و نشانه‌های متداول زبان انگلیسی هستند، محدود می‌کنند. مشکل بخش‌های دودویی آن است که اغلب در سطح بیت تفسیر می‌شوند و یادگیری وابستگی‌های آنها با شبکه‌های عصبی، که داده‌های محدودی را در زمان آموزش می‌بینند، امکان‌پذیر نیست. از طرفی هنگامی که نسبت کوچکی از داده‌ها به صورت دودویی در میان قسمت‌های متنی ظاهر می‌شوند می‌توان آنها را با روش‌های جابه‌جایی تصادفی فاز کرد.

ما در روش پیشنهادی خود داده‌های دودویی را از قالب فایل حذف کرده اما به جای آن یک توکن خاص موسوم به توکن دودویی^۳ را جایگزین می‌کنیم. بدین ترتیب در فرایند آموزش مدل توزیع آماری نحوه ظاهر شدن قسمت‌های دودویی نیز فراگیری می‌شود. حال در فرایند تولید داده‌های آزمون جدید مدل مکان‌های احتمالی ظاهر شدن قسمت دودویی را با پیش‌بینی وقوع توکن مربوطه، تعیین می‌کند. سپس توکن را با یک قسمت دودویی که به صورت تصادفی (یا هر روش دیگر) فاز شده است جایگزین می‌کنیم. یعنی وارون عملی که پیش از فرایند آموزش انجام دادیم را پس از فرایند تولید، انجام می‌دهیم. با کلیت گفته شده در اینجا یک روش ترکیبی تولید خودکار داده آزمون خواهیم داشت.

^۱Corpus

^۲American Standard Code for Information Interchange

^۳Binary Token



شکل ۴-۱: روندنمای روش پیشنهادی در حالت کلی.

۴-۱-۲ تمایز داده و فراداده

در فصل ۱، تمایز داده و فراداده به عنوان یک مسئله در تولید داده آزمون مطرح گردید. مدل مطرح در روش پیشنهادی این امکان را فراهم می‌کند تا بتوانیم به صورت احتمالی داده و فراداده را از یکدیگر تشخیص دهیم. از آنجایی که فراداده‌ها در ساختار یک فایل از پیش تعریف شده و تقریباً ثابت هستند، تکرار آنها به مراتب بیشتر از تکرار داده‌ها است که محتویات یک فایل را تشکیل می‌دهند. مدل بدین ترتیب در توزیع آماری یادگیری شده احتمال بیشتری را به فراداده‌ها نسبت می‌دهد. با تعیین یک آستانه مشخص برای هریک از انواع داده و فراداده در هنگام تولید داده آزمون قادر به تصمیم‌گیری در مورد نوع آنها هستیم. ما نتیجه این تصمیم را در نحوه فاز داده آزمون جدید به کار می‌بریم، بدین صورت که دو الگوریتم فاز با اهداف فاز فراداده برای مرحله تجزیه و فاز داده برای مرحله پرداخت فایل ارائه خواهیم داد.

۴-۱-۳ مثال انگیزشی

تمرکز اصلی ما در روش پیشنهادی، بر روی پیمانه تولید داده آزمون در یک فازر قالب فایل است. شکل ۴-۲، مراحل تولید داده آزمون از طریق روش پیشنهادی را روی قالب فایل PDF، نشان می‌دهد. سه مرحله اصلی در این شکل وجود دارد که به ترتیب با شماره‌های یک تا سه، شماره‌گذاری شده‌اند. مطابق این مراحل، ابتدا پیکره‌ای از فایل‌های از پیش موجود و معتبر جمع آوری گردیده و به عنوان مجموعه داده انتخاب می‌شوند. در مرحله اول، یعنی پیش‌پردازش، هریک از فایل‌های ورودی پیمایش شده، بخش‌های متنی آن جدا و به یکدیگر الحاق می‌شوند. همچنین بخش‌های دودویی داخل فایل، نیز به صورت مجزا، پس از جدا شدن، نگهداری می‌شوند. در فایل PDF، بخش‌هایی دودویی بین کلیدواژه‌های stream و endstream می‌آیند که در پیش‌پردازش آنها را با توکن دودویی stream جایگزین می‌کنیم. خروجی این مرحله دو مجموعه است. یک مجموعه داده متنی و یک مجموعه داده دودویی. در مجموعه داده متنی، مکان قسمت‌های دودویی با توکن دودویی مشخص شده است.

در مرحله دوم، یک مدل زبانی عصبی بر روی مجموعه داده حاصل از جداسازی بخش‌های متنی پیکره آموزش داده می‌شود. مانند هر وظیفه یادگیری ماشینی در این مرحله ابتدا مجموعه داده ورودی به سه مجموعه آموزش، آزمون و ارزیابی تقسیم می‌گردد. خروجی این مرحله یک مدل مولد است که قادر به تولید داده‌های جدید با پیروی از توزیع حاکم بر داده‌های مجموعه آموزش بوده و کاراکتر به کاراکتر داده‌های جدید را تولید می‌کند.

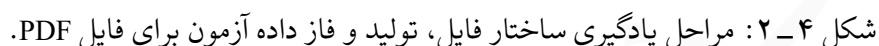
مرحله سوم با روش‌هایی اقدام به تولید داده‌های آزمون جدید، در این جا فایل‌های PDF، می‌کند. در این مرحله بخش‌های متنی فایل جدید از طریق مدل مولد، یعنی با روش مبتنی بر تولید، ایجاد می‌گردد. عمل

بدشکل کردن داده‌های متنی نیز همزمان با تولید آنها صورت می‌گیرد. بدین صورت که برخی از کاراکترها بعد از تولید، بر اساس شرایطی که در ادامه فصل توضیح داده می‌شوند، با کاراکترهای دیگری جایگزین می‌شوند. در همین مرحله، بخش‌هایی دودویی نیز در صورت پیش‌بینی توسط مدل مولد، با یک روش مبتنی بر جابه‌جایی، تولید و در مکان مناسب خود قرار می‌گیرند. منظور از پیش‌بینی در اینجا، تولید توکن دودویی stream در توسط مدل مولد است. در تولید بخش دودویی یک داده آزمون از بخش‌های دودویی پیکره اولیه که در مرحله یک جداسازی شده است به‌عنوان دانه اولیه استفاده می‌کنیم. در واقع آنها را با یکی از عملگرهای جابه‌جایی که در فصل ۲ به آن اشاره شد به‌سادگی جابه‌جا کرده و یک داده آزمون دودویی جدید و فاز شده می‌سازیم.

مرحله سوم در مجموع تعدادی داده آزمون را با روشی ترکیبی، تولید می‌کند. این داده‌های آزمون سپس می‌توانند به عنوان داده‌های آزمون در فرایند آزمون فازی استفاده شوند. همچنین می‌توان از آن به عنوان دانه اولیه برای فازهای قالب فایل مبتنی بر جابه‌جایی نیز استفاده کرد. برای ارزیابی روش پیشنهادی خود که در فصل بعد به آن خواهیم پرداخت، در انتهای این فصل یک فازر قالب فایل ساده را مطرح می‌کنیم. این فازر قادر است تا پس از تولید داده‌های آزمون به روش شرح داده شده آنها را به عنوان ورودی به SUT تزریق کرده و خطاهای احتمالی را در صورت وجود گزارش کند. مرحله پیش‌پردازش بسیار ساده بوده و نیاز به توضیح خاصی ندارد. این مرحله ممکن است برای قالب‌های فایل مختلف، کاملاً متفاوت باشد. مثلاً برای قالب‌های فایل متنی مثل HTML و XML پیش‌پردازش تنها الحاق کردن کلیه فایل‌های موجود در پیکره است. مراحل دوم و سوم، یعنی ایجاد و آموزش مدل مولد و تولید و فاز داده آزمون، به ترتیب در بخش ۴-۲ و بخش ۴-۳ توضیح داده شده‌اند. همچنین معماری فازر قالب فایل پیشنهادی در بخش ۴-۴ بیان گردیده است.

۴-۲ مدل

یک فایل را می‌توان نمونه تولید شده توسط زبانی که قالب آن فایل را بیان می‌کند دانست. با این فرض می‌توان یک مدل زبانی برای هر قالب فایلی ایجاد کرد. مدل زبانی عصبی عملکرد بهتری نسبت به مدهای زبانی سنتی دارد. با توجه به ذات مبتنی بر توالی بایت‌های یک فایل، وابستگی بایت فعلی به یک توالی از بایت‌های قبلی و نیز تجزیه ترتیبی (چپ به راست و بالا به پایین) آن توسط تجزیه‌گر SUT در بسیاری از موارد، در روش پیشنهادی خود، از RNN برای ایجاد مدل زبانی و یادگیری ساختار فایل استفاده می‌کنیم. در فصل ۲ دیدیم که RNN برای یادگیری وظایف مبتنی بر توالی ابداع شده است. به‌طور دقیق‌تر در هنگام آموزش از یک مدل RNN چند به یک، با معماری شکل ۲-۶ (پ) استفاده خواهیم کرد و در هنگام تولید داده جدید با فراخوانی همین مدل به‌صورت متوالی، مدلی از نوع شکل ۲-۶ (ت) خواهیم داشت.



LSTM دوسویه توالی ورودی را به دو صورت روبه جلو^۳ و روبه عقب^۴ می بیند. در واقع LSTM دوسویه از دو LSTM یکسویه تشکیل شده است. یکی از آنها توالی را از چپ به راست پردازش می کند و دیگری از راست به چپ. در نتیجه هر گذرجلو LSTM دو خروجی خواهد داشت. یک تابع ادغام^۵ برای ادغام خروجی هریک از LSTM های روبه جلو و روبه عقب به کار می رود. تابع ادغام می تواند هر تابعی که قادر به ترکیب دو بردار با طول های یکسان است، در نظر گرفته شود؛ از جمله: جمع، ضرب، الحاق و غیره. ما از تابع ادغام جمع برای مدل پیشنهادی خود استفاده می کنیم که درایه های نظیر به نظیر هر یک از بردارهای خروجی را با یکدیگر جمع می زند.

علاوه بر نوع مدل LSTM، برای LSTM یک سوپیه مدل هایی با تعداد عصب های متفاوت و نیز تعداد

¹Unidirectional²Bidirectional³Forward⁴Backward⁵Merge Function

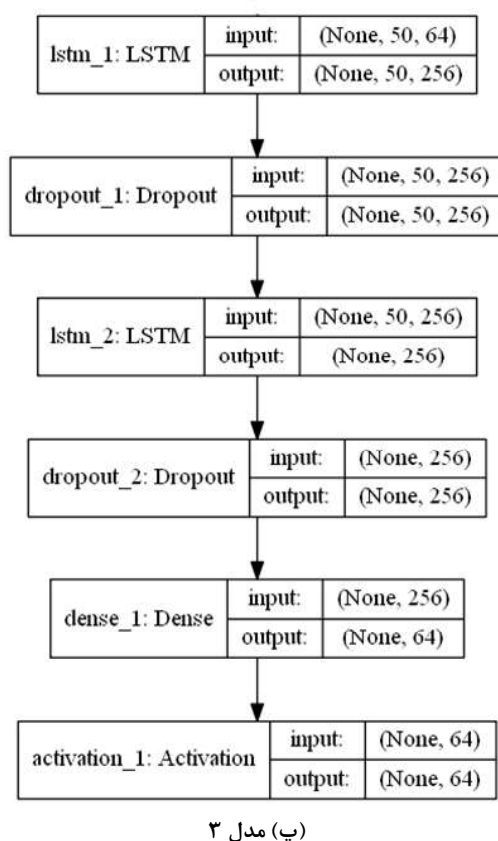
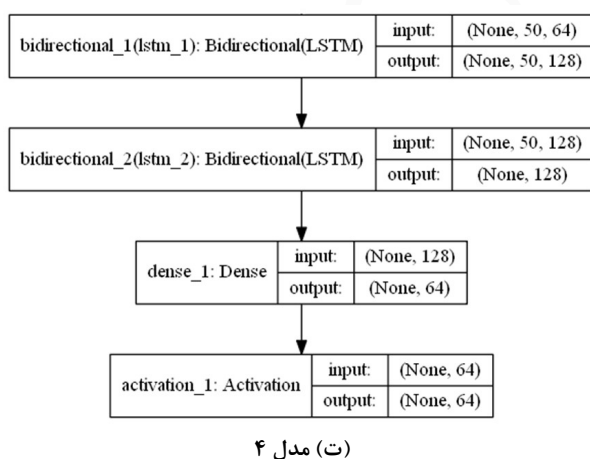
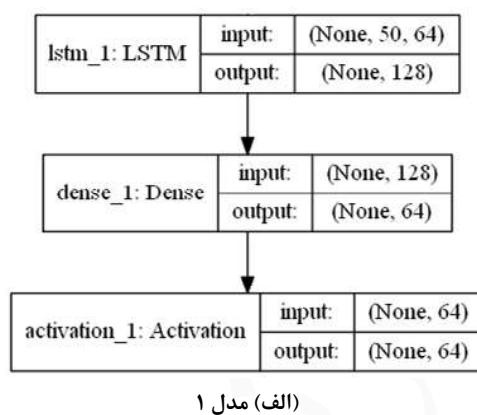
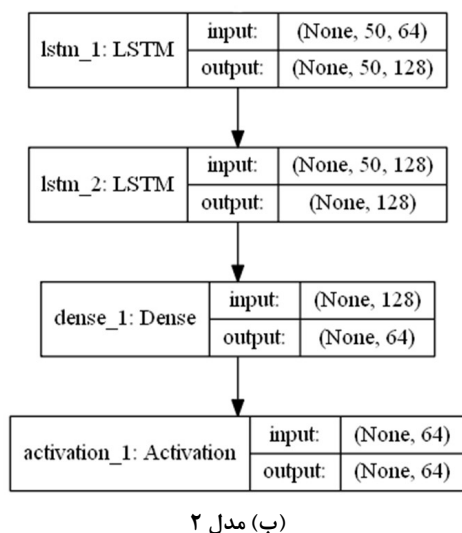
جدول ۴-۱: مدل‌های یادگیری ژرف طراحی شده، پارامترها و ابرپارامترهای آنها.

شماره (نام) مدل	نوع (معماری) مدل	تعداد لایه پنهان	تعداد عصب در هر لایه	تعداد پارامترها
۱	Unidirectional LSTM (Many to One)	۱	۱۲۸	۱۲۷۵۸۴
۲	Unidirectional LSTM (Many to One)	۲	۱۲۸	۲۳۸۶۵۶
۳	Unidirectional LSTM (Many to One)	۲	۲۵۶	۸۷۰۴۶۴
۴	Bidirectional LSTM (Many to One)	۲	۱۲۸	۴۶۹۰۵۶

لایه‌های پنهان متفاوت را ساخته‌ایم. هدف از این کار مشاهده تأثیر مدل‌های با پیچیدگی مختلف در یادگیری ساختار فایل و تولید داده‌های آزمون و آزمایش این فرضیه است که آیا مدل‌های پیچیده‌تر، مدل‌های ساده‌تر را شکست می‌دهند؟ یعنی موفق به افزایش پوشش کد و یا شناسایی خطاهای بیشتری می‌شوند یا خیر. جدول ۴-۱ مدل‌های طراحی شده برای استفاده در آزمایش‌های روش پیشنهادی و مشخصه‌های هریک را نشان می‌دهد. یکی از مهم‌ترین مشخصه‌ها تعداد پارامترهای هر مدل است. گراف محاسباتی هریک از مدل‌های جدول ۴-۱ در شکل‌های ۴-۳ الف تا ۴-۳ ت آمده است. اندازه ماتریس‌های ورودی و خروجی در گراف مرتبط با تعداد لایه‌های پنهان و نیز تعداد عصب‌ها در هر لایه است که در ادامه بیشتر توضیح داده خواهد شد. همچنین در ادامه پایان‌نامه، هریک از مدل‌های معرفی شده، با شماره ذکر شده برای آنها در جدول ۴-۱ شناخته می‌شوند.

۴-۲-۱ آموزش مدل

فرایند آموزش همه مدل‌های جدول ۴-۱ یکسان است. در واقع برای آموزش هر مدل تنها نیاز داریم که ورودی و خروجی شبکه عصبی ژرف را مشخص نماییم. پس از استخراج مجموعه داده اولیه که حاوی یک پیکره متوالی از کاراکترهای در بردارنده فایل است، آن را به مجموعه‌های آموزش، آزمون و ارزیابی تقسیم می‌کنیم. چون خروجی مدل مولد در هربار اجرا یک توالی خواهد بود، طول توالی‌های مجموعه داده معیار مناسبی برای تقسیم‌بندی آن به مجموعه‌های سه‌گانه یاد شده است به نحوی که هر مجموعه توزیع یکسانی از فراوانی طول‌های مختلف باشد. از مجموعه‌های آموزش و ارزیابی در هنگام آموزش مدل و از مجموعه آزمون در هنگام تولید داده‌های جدید از روی مدل، استفاده خواهیم کرد. تفکیک مجموعه داده به سه مجموعه



شکل ۴-۳: گراف محاسباتی مدل‌های عصبی ژرف جدول ۴-۱. اعداد داخل مستطیل‌ها اندازه ماتریس‌های ورودی و خروجی هر لایه را نشان می‌دهند. گراف‌ها با استفاده از توابع کتابخانه Keras [۵۰] تولید شده‌اند.

آموزش، ارزیابی و آزمون در وظایف یادگیری ماشینی یک امر بدیهی است، اما استفاده از آن برای آزمون فازی و تولید داده‌های آزمون جدید تا قبل از این، انجام نشده و چنانچه در ادامه خواهیم دید، ایده جدیدی است. در واقع ما از داده‌های مجموعه آزمون برای تولید داده‌های آزمون جدید استفاده خواهیم کرد.

هر مجموعه بدین ترتیب حاوی تعدادی فایل است و هر فایل یک توالی از کاراکترهای ASCII است. شروع و پایان هر یک از این توالی‌ها پیش از فرایند آموزش، با توکن‌های شروع و پایان که نشانه‌های مجزایی (بیرون از مجموعه داده) هستند برچسب‌گذاری می‌شود. البته بسیاری از قالب‌های فایل چنین توکن‌هایی را به عنوان بخشی از ساختار خود دارند. برای مثال فایل‌های HTML با برچسب `<html>` شروع و با برچسب `</html>` خاتمه می‌یابند. فایل‌های XML و PHP نیز چنین هستند. در این حالت نیازی به افزودن توکن‌های شروع و پایان جدید نیست و می‌توان از همین توکن‌ها استفاده کرد. دلیل افزودن توکن‌های شروع و پایان در این مرحله، به نحوه ساخت مدل‌های زبانی عصبی باز می‌گردد. در واقع این مدل‌ها با دیدن توکن آغازین پیش‌بینی را شروع و با تولید توکن پایانی، پیش‌بینی را خاتمه می‌دهند.

به منظور آموزش هر مدل چند به یک و ایجاد یک مدل زبانی عصبی، ابتدا با الحاق محتوای فایل‌های مجموعه آموزش، یک توالی بزرگ از کاراکترها به شکل $S = \langle s^{(1)}, s^{(2)}, s^{(3)}, \dots, s^{(n)} \rangle$ ایجاد می‌کنیم. سپس این توالی را به تعدادی توالی ورودی کوچکتر با طول ثابت d می‌شکنیم طوری که i امین توالی ورودی برابر است با:

$$x_i = S[i * j : (i * j) + d] \quad (4-1)$$

در رابطه ۴-۱، $s[l, u]$ برشی از توالی s بین شاخص‌های l و u بوده و j میزان پرش به جلو در انتخاب توالی ورودی بعدی از توالی اصلی را نشان می‌دهد. ابرپارامتر d در اینجا، در واقع همان بردار زمینه در مدل زبانی است. انتخاب d به پارامترهایی مانند فاصله وابستگی‌های بین بایت‌های مختلف در یک فایل و طول فایل‌های مجموعه داده، بستگی دارد. وقتی از LSTM به عنوان بلوک‌های سازنده RNN در ساخت NLM استفاده شود، می‌توان وابستگی‌های طولانی‌تر را به خوبی یاد گرفت. در عمل و برای وظایف مدل زبانی، این مقدار را بین ۴۰ تا ۱۰۰ نشانه در نظر می‌گیرند.

گفتیم شبکه عصبی در حالت بانظارت آموزش می‌بیند؛ یعنی، برای هر ورودی یک برچسب خروجی نیاز است. در اینجا کاراکتر بعدی، که هدف پیش‌بینی آن است، را به عنوان برچسب خروجی برای هر توالی ورودی نسبت می‌دهیم. بنابراین برچسب خروجی برای i امین توالی ورودی برابر خواهد بود با:

$$Y_{x_i} = S[(i * j) + d + 1] \quad (4-2)$$

پس از تولید همه نمونه‌های آموزشی و برچسب‌های متناظر آنها، مدل به صورت انتها به انتها آموزش داده می‌شود؛ مشابه روش یادگیری و فاز [۱۱]. یعنی، داده‌های خام مجموعه آموزش، بدون استخراج هیچ‌گونه ویژگی خاصی به شبکه داده می‌شوند و در زمان استفاده نیز، شبکه داده‌های نهایی را بدون نیاز به هیچ‌گونه پس‌پردازشی تولید می‌کند. در این حالت مدل قادر پیش‌بینی مقدار احتمال شرطی $p(x^{(i+d+1)} | x^{(i)}, \dots, x^{(i+d)}) >$ خواهد بود. روند تولید نمونه‌های آموزشی در الگوریتم ۴-۱ نشان داده شده است. همچنین شکل ۴-۴، یک فایل HTML و سه توالی آموزشی اول تولید شده برای آن توسط الگوریتم ۴-۱ را به همراه موقعیت قرارگیری آنها در شبکه نشان می‌دهد. توالی ورودی کلی در این مثال، همه کاراکترهای فایل HTML و پارامترهای d و j به ترتیب برابر ۳ و ۱ قرار داده شده‌اند. مقادیر در نظر گرفته شده کاملاً فرضی هستند.

شبکه عصبی با مقادیر عددی کار می‌کند. بنابراین بایستی یک نمایش عددی برای هر کاراکتر در مجموعه آموزش در نظر گرفت. یک روش مرسوم در مدل‌های مبتنی بر کاراکتر استفاده از بردارسازی One-hot یا ۱ از k است. در این روش یک بردار به طول اندازه مجموعه واژگان برای هر کاراکتر در نظر می‌گیرند که تنها یکی از شاخص‌های آن برابر ۱ است. بدین ترتیب برای نمایش V کاراکتر، V بردار متفاوت خواهیم داشت. در اینجا نیز از همین روش استفاده می‌کنیم.

خروجی شبکه نیز که به شکل یک کاراکتر تعریف شده است در واقع یک بردار One-hot است. در فرایند آموزش تنها یکی از شاخص‌های این بردار یک است. اما در مرحله آزمون بردار تولیدی حاوی مقادیر حقیقی مختلفی در بازه $(-\infty, +\infty)$ برای هر شاخص است. با اعمال تابع بیشینه هموار (رجوع شود به رابطه ۲-۱۲)، در خروجی این لایه، می‌توان خروجی را به یک توزیع احتمالی معتبر تبدیل کرد.

الگوریتم ۴-۱ ProduceTrainingSamples

Input: Total sequence S , Input sequences length d , Jump step j

Output: Input sequences x , Output labels Y

```

1  $x \leftarrow \text{List}()$  /* Make an empty list */
2  $Y \leftarrow \text{List}()$  /* Make an empty list */
3 for  $i \leftarrow 1$  to  $\text{Len}(S) - d$  do
4    $x[i] \leftarrow S[i * j : (i * j) + d]$ 
5    $Y[i] \leftarrow S[(i * j) + d + 1]$ 
6 end
7 Return  $x, Y$ 
```

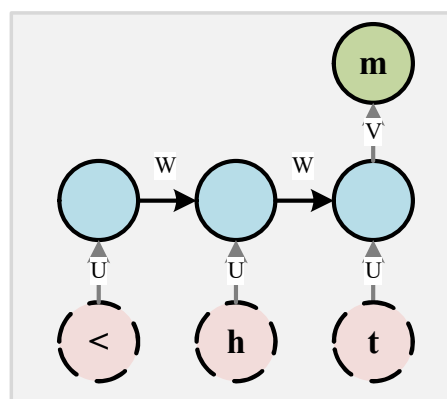


```

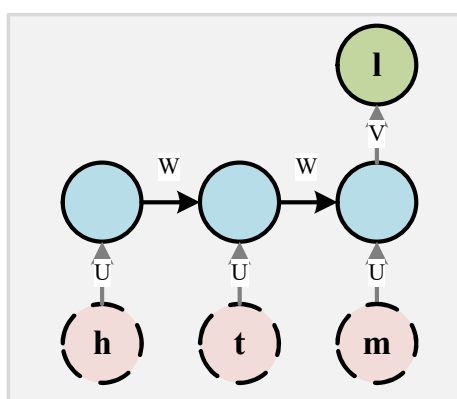
<html>
<body>
<h1>Data</h1>
</body>
</html>

```

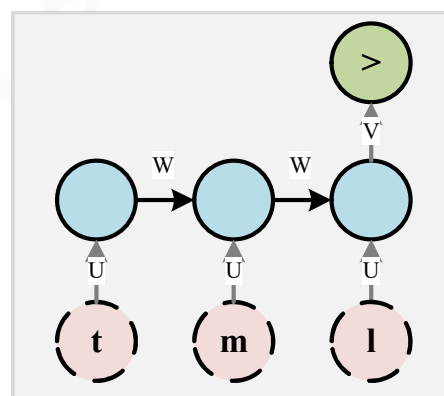
(الف) یک فایل HTML



(ب) توالی آموزشی اول



(پ) توالی آموزشی دوم



(ت) توالی آموزشی سوم

شکل ۴-۴: مثالی از نحوه آموزش مدل‌های چند به یک روی یک فایل HTML ساده. تنها سه توالی آموزشی اول در این شکل نشان داده شده است. اما الگوریتم ۴-۱ کل فایل را به توالی‌های آموزشی تقسیم می‌کند.

۴-۲-۲ تولید داده‌های جدید

پس از اتمام فرایند آموزش، مدل برای تولید داده‌های جدید مورد پرس‌وجو قرار می‌گیرد. در این مرحله ابتدا یک پیشوند ثابت از کاراکترهای شروع کننده یک توالی از فایل‌های مجموعه آزمون به طول d ، که توکن شروع در ابتدای آن است، به شبکه خورانیده می‌شود^۱. شبکه توزیع احتمالی کاراکتر $d+1$ ام را به‌عنوان خروجی تولید می‌کند که شامل یک بردار One-hot از مقادیر احتمال‌های تمام کاراکترهای دیده شده هنگام آموزش است. سپس یک کاراکتر از این توزیع احتمالی، انتخاب شده و به پیشوند ورودی الحاق می‌گردد. در این حالت طول پیشوند $d+1$ است. در مرحله بعد سمت چپ‌ترین کاراکتر پیشوند حذف شده و d کاراکتر باقی مانده دومرتبه به شبکه داده می‌شود تا توزیع احتمالی کاراکتر $d+2$ حاصل شود. این روند تا زمان تولید توکن پایانی که خاتمه یافتن یک توالی را پیش‌بینی می‌کند، ادامه خواهد داشت. در [۱۱] سه راهبرد تولید داده جدید از توزیع احتمالی پیش‌بینی شده توسط مدل معرفی شده که در بخش ۳-۳-۲ آنها را دیدیم.

در اینجا ما از راهبرد نمونه‌برداری که روشی مرسوم برای تولید داده از یک توزیع آماری است استفاده می‌کنیم. یعنی هر بار از توزیع احتمالی پیش‌بینی شده به عنوان یک توزیع چندجمله‌ای^۲ نمونه‌برداری می‌نماییم. البته راهبردهای جدیدی را نیز معرفی خواهیم کرد. در زبان پایتون و با استفاده از بسته numpy می‌توان به‌صورت زیر از یک توزیع چندجمله‌ای نمونه‌برداری کرد:

```

1 import numpy as np
2 sample = np.random.multinomial(n, prob_vals, size=None)

```

برنامه ۴-۱: نمونه‌برداری از توزیع چندجمله‌ای در پایتون.

پارامتر تنوع

در [۱۱] انتخاب حریصانه به دلیل وجود پیشوند ثابت راهبرد ناکارمندی معرفی شده است. در روش پیشنهادی ما، مدل پیشوندهای متغیری را می‌پذیرد. بنابراین می‌توانیم از راهبرد حریصانه هم استفاده نماییم. با این حال تعداد داده‌های تولیدی در این راهبرد محدود به تعداد نمونه‌های مجموعه آزمون است که البته تعداد کمی نخواهند بود.

^۱ دقت شود که ویژگی فایل‌های مجموعه آزمون آن است که قبلاً در فرایند آموزش توسط مدل مشاهده نشده‌اند و این امر تأثیر مستقیمی در ارزیابی خوب بودن مدل و تنوع داده‌های تولید شده توسط مدل دارد. بدون وجود مجموعه آزمون نمی‌توانیم معیارهایی مانند سرگشتگی را به‌درستی حساب کنیم. به همین سبب جداسازی مجموعه آزمون را تأکید می‌کنیم.

^۲Multinomial Distribution

در مقابل راهبرد حریصانه، نمونه‌برداری باعث ایجاد اشیای متنوع می‌شود که لزوماً خوش‌شکل نیستند. به همین دلیل پارامتری به نام تنوع^۱ را با هدف کنترل داشتن بر روی تنوع داده‌های تولیدی مطرح می‌کنیم. تنوع D به صورت یک عدد حقیقی در بازه $(0, +\infty)$ تعریف می‌گردد. در هنگام آزمون مدل، مقادیر حقیقی پیش‌بینی شده توسط مدل، ابتدا بر D تقسیم می‌شوند و سپس تابع بیشینه همواره به آن اعمال می‌شود تا بردار خروجی حاصل گردد. در نتیجه هرچه قدر مقدار D کوچک‌تر (نزدیک به صفر) انتخاب شود، نمونه‌برداری به راهبرد حریصانه نزدیک شده، داده‌های خوش‌شکل‌تری تولید می‌گردد اما از تنوع آن کاسته می‌شود. بالعکس هرچه قدر مقدار D بزرگتر از یک انتخاب شود، اختلاف بین مقادیر پیش‌بینی شده کم، تنوع داده‌های تولید شده زیاد و متقابلاً احتمال بدشکل بودن آنها افزایش می‌یابد.

n- فهرست بهتر

یک راهبرد دیگر تولید داده جدید، استفاده از n- فهرست بهتر است. در این راهبرد در هر بار پیش‌بینی مدل پس از اعمال تابع بیشینه هموار، تعداد n احتمال بالای بردار خروجی انتخاب و نگهداری می‌شود. سپس با هریک از این احتمال‌ها پیش‌بینی مرحله بعد انجام می‌شود. برای انتخاب n- فهرست بهتر بین دو پیش‌بینی مقادیر احتمال‌های هر دو فهرست در یکدیگر ضرب و n حاصل ضرب بالاتر انتخاب می‌شود. پس از B مرحله، n پیشوند B تایی خواهیم داشت که پیشوند با بالاترین احتمال پیشوند خروجی خواهد بود. این راهبرد البته زمان و حافظه مصرفی بالایی نیاز دارد. لذا مقادیر n و B بایستی کوچک انتخاب شوند. در ترجمه ماشینی معمولاً از این راهبرد نیز استفاده می‌شود. ما در این پایان‌نامه و در آزمایش‌های خود، راهبرد بالا را ارزیابی نکردیم، اما می‌توان از آن در عمل استفاده کرد.

۴-۳ الگوریتم‌های فاز عصبی

این بخش را می‌توان هسته روش پیشنهادی و نوآوری‌های ارایه شده در این پایان‌نامه دانست. در اینجا توضیح می‌دهیم که داده‌های آزمون چگونه با روشی خودکار و ترکیبی تولید و فاز می‌شوند. هنگامی که داده‌های آزمون را با استفاده از مدل‌های بخش ۴-۲ و راهبردهای معرفی شده در بخش ۴-۲-۲ تولید می‌کنیم یک تنوع ذاتی در داده‌ها وجود خواهد داشت و در هر صورت داده‌های آزمون جدید محسوب می‌شوند. اما اهداف یادگیری قالب فایل و آزمون فازی قالب فایل در تضاد بایکدیگر هستند. هدف الگوریتم یادگیری ایجاد داده‌های خوش‌شکل است که بتواند از تجزیه‌گر اولیه عبور و مسیرهای عمیق را اجرا کند و هدف آزمون فازی بدشکل

¹Diversity

کردن ورودی است به نحوی که تجزیه‌گر و دیگر قسمت‌های کد اجرا شده را به حالت خرابی ببرد. در این بخش، ما الگوریتم‌های فاز عصبی را با هدف ایجاد یک مصالحه بین دو هدف قبلی و تولید نهایی داده آزمون برای یک فازر قالب فایل، معرفی می‌کنیم.

فاز کردن داده آزمون همزمان با تولید آن از روی مدل صورت می‌گیرد و مرحله مجزایی نیست. الگوریتم‌های ۲-۴ و ۳-۴ نحوه تولید داده آزمون که در بخش ۲-۲-۴ بحث شد و فاز همزمان آن را نشان می‌دهند. هر دو الگوریتم به عنوان ورودی مدل یادگیری شده M ، پیشوند شروع توالی P (حاوی توکن شروع)، تنوع نمونه‌برداری D ، نرخ فاز FR ، توکن پایان ET و توکن دودویی BT را پذیرفته و داده آزمون TD را به عنوان خروجی باز می‌گردانند.

حلقه اصلی هر یک از الگوریتم‌ها، تا زمانی که ET تولید نشده باشد ادامه می‌یابد. چنانچه طول توالی تولید شده از یک حد آستانه $MaxLen$ بیشتر شد ولی ET کماکان تولید نشده باشد، آنگاه ET به صورت خودکار به انتهای TD اضافه می‌شود تا الگوریتم از حلقه تکرار بیرون آید. بدین ترتیب مطمئن می‌شویم که الگوریتم همواره خاتمه خواهد یافت. سپس در صورتی که مدل حضور یک داده دودویی در TD را پیش‌بینی کرده باشد (یعنی BT در توالی TD پیدا شود)، ابتدا یک بخش دودویی جایگزین BT شده و سپس این بخش با روش‌های جابه‌جایی تصادفی یا هر روش دلخواه دیگر به صورت مجزا فاز می‌شود و در نهایت داده آزمون نهایی توسط الگوریتم بازگردانیده می‌شود. ما در عمل برای خط ۱۹ هر دو الگوریتم، عملگر معکوس کردن یک بایت را در نظر گرفتیم. بدین ترتیب که درصد ثابتی از بایت‌های دودویی را در هر بخش دودویی انتخاب شده، به صورت تصادفی انتخاب نموده و مقادیر بیت‌های هر بایت انتخابی را معکوس می‌کنیم.

۴-۳-۱ فاز عصبی داده

تفاوت اصلی در الگوریتم‌های فاز عصبی داده و فاز عصبی فراداده در هنگام اعمال عملیات فاز است. در بخش ۲-۱-۴ گفتیم مدل به داده‌ها احتمال کمتری نسبت می‌دهد. بنابراین در الگوریتم ۲-۴، هرگاه احتمال کاراکتر پیش‌بینی شده از یک حد آستانه α کمتر باشد، متوجه می‌شویم که کاراکتر جزو بخش داده فایل بوده که در این حالت آن را با کاراکتری که کمترین احتمال را دارد جایگزین می‌کنیم. با این امید که کم‌احتمال‌ترین کاراکتر، ناخواسته‌ترین کاراکتر نیز بوده و ممکن است منجر به خرابی شود.

خط ۷ الگوریتم ۲-۴ بیان می‌دارد که برای فاز بایستی شرایط دیگری علاوه بر شرط فوق برقرار باشد. از جمله عدد تصادفی p_fuzz که لازم است از FR یعنی نرخ فاز کمتر شود. هرچه قدر مقدار FR بزرگتر انتخاب شود شرط $p_fuzz < FR$ برای اعداد بیشتری تحقق می‌یابد. بدین ترتیب آزمون‌گر می‌تواند روی درصد کاراکترهای فاز شده کنترل داشته باشد. همچنین می‌خواهیم که مجموعه کاراکترهای توالی‌های ET و

BT فاز نشود؛ زیرا از آنها در قسمت‌های بعد استفاده می‌کنیم. لذا دو شرط دیگر به شرط‌های خط ۷ الگوریتم ۴-۲ اضافه می‌شوند. بنابراین فاز کاراکتر پیش‌بینی شده تنها در حالتی که ترکیب عطفی هر ۴ شرط درست باشد، صورت می‌گیرد.

تغییر یک بایت یا یک کاراکتر از بخش‌های داده یک فایل، معمولاً بدشکل بودن خاصی را ارائه نمی‌دهد. به طور شهودی تصور کنید که به جای عدد ۱۲۵ عدد ۹۲۵ قرار داده شود یا هر البته هر کاراکتر دیگری در موقعیت رقم ۱. در بد شکل کردن داده هدف ما جایگزین کردن یک مقدار مرزی به جای داده موجود است. مثلاً برای مورد گفته شده جایگذاری عددی به فرم ۹۹۹...۹ خوب به نظر می‌رسد. در واقع در فاز عصبی داده، به دنبال روشی برای انتشار فاز به کاراکترهای بعدی نیز هستیم. وقتی یک کاراکتر فاز شد، دو راهکار برای تولید کاراکترهای بعدی پیش‌رو داریم. راهکار اول تولید کاراکترهای بعدی، مستقل از کاراکتر فاز شده است؛ یعنی اینکه مدل در تولید کاراکتر t ام هیچ آگاهی از فاز شدن کاراکتر مرحله $t-1$ ام خود نداشته باشد. برای این منظور کافی است تا کاراکتر اولیه تولید شده توسط مدل، یعنی کاراکتر قبل از فاز و جایگزین شدن، را به پیشوند شروع توالی P اضافه کنیم. بدین ترتیب مدل مستقل از اینکه کاراکتر فاز شده است یا نه به تولید داده‌های مبتنی بر گرامر (داده‌های خوش‌شکل) ادامه می‌دهد. راهکار دوم تولید کاراکترهای بعدی، وابسته به کاراکتر فاز شده است؛ یعنی اینکه اثر کاراکتر فاز شده به کاراکترهای تولید شده در مراحل بعد سرایت کند. برای این منظور کاراکتر فاز شده را به پیشوند شروع توالی P اضافه کرده و مدل را به سمت بد شکل کردن کاراکترهای بعدی متمایل می‌کنیم. این راهکار در الگوریتم فاز عصبی داده به دلیل شرح داده شده استفاده گردیده است. خط ۱۰ الگوریتم ۴-۲، کاراکتر فاز شده یعنی کاراکتر c را به پیشوند P اضافه می‌کند. در همین خط برای ثابت ماندن طول پیشوند که در واقع ورودی مدل است، کاراکتر اول یعنی قدیمی‌ترین کاراکتر را از پیشوند حذف می‌کنیم. در کل این خط عمل انتشار فاز به جلو را پیاده‌سازی می‌کند.

۴-۳-۲ فاز عصبی فراداده

در الگوریتم فاز عصبی فراداده (الگوریتم ۴-۳)، می‌خواهیم کاراکترهای فراداده را تغییر دهیم. برای این منظور بررسی می‌کنیم که احتمال کاراکتر پیش‌بینی شده توسط مدل یعنی $p(c)$ بیشتر از حد آستانه β باشد، در این صورت کاراکتر با کمترین احتمال را جایگزین می‌کنیم. در اینجا اما به دنبال تغییر تنها یکی از کاراکترها در هر توکن فراداده هستیم. لذا در خط ۱۱ الگوریتم ۴-۳ تغییری در پیشوند بعدی مدل ایجاد نمی‌کنیم تا مدل به همان روند تولید داده خوش‌شکل خود ادامه دهد بدون اینکه متوجه فاز شدن کاراکتر پیش‌بینی کرده گام زمانی قبلی خود باشد. در واقع در فاز عصبی فراداده، انتشار فاز به جلو را انجام نمی‌دهیم. ایده ما برای این کار این است که در اغلب مواقع یک تغییر کوچک، مثلاً تغییر یک بایت در قسمتی که قالب فایل را بیان

می‌کند سبب گمراه شدن تجزیه‌گر می‌شود. این عمل درست در خلاف مورد داده است که یک بایت معمولاً تغییری ایجاد نمی‌کند و این مقادیر مرزی (بیش از حد کوچک، بیش از حد بزرگ، خالی یا صفر) هستند که مسبب خرابی می‌شوند. مقادیر داده معمولاً به صورت یک توالی از بایت‌ها ادامه می‌یابند.

در الگوریتم ۳-۴ پارامتر نرخ فاز FR همچنان وجود دارد، اما دو شرطی که بررسی کننده وجود توکن‌های BT و ET هستند را حذف کردیم تا علاوه بر آسان‌تر کردن شرط فاز بتوانیم آنها را نیز به عنوان بخشی از قالب فایل، فاز کنیم؛ زیرا، تعداد خطاهایی که در تجزیه‌گر رخ می‌دهند، معمولاً بیشتر هستند. در نهایت خطوط سایه زده شده در الگوریتم ۳-۴ (خطوط ۷، ۸ و ۱۱) محل‌های تفاوت این الگوریتم با الگوریتم فاز عصبی داده (الگوریتم ۲-۴) را نشان می‌دهد.

حد آستانه $MexLen$ که حداکثر طول داده آزمون تولیدی را در الگوریتم‌های ۲-۴ و ۳-۴ مشخص می‌کند نیز به صورت یک عدد صحیح تصادفی در بازه $[a, b]$ مشخص می‌شود. مقادیر این بازه می‌تواند براساس میانگین طول فایل‌های مجموعه آزمون، انتخاب شود. بدین صورت که کران‌های آن برابر اختلاف واریانس طول از میانگین باشند. مقادیر ابرپارامترهای موجود در هر دو الگوریتم بایستی توسط آزمون‌گر مطابق قالب فایل مورد آزمون و نکته‌های بیان شده، تنظیم گردد. در فصل ۵ که ارزیابی روش پیشنهادی را مطرح می‌کنیم، مقادیر در نظر گرفته شده را برای قالب فایل PDF بیان می‌کنیم.

۴-۴ پیاده‌سازی

جزئیات پیاده‌سازی روش پیشنهادی در پیوست ب آمده است. برای پیاده‌سازی مدل‌های یادگیری ژرف از کتابخانه سطح بالای یادگیری ژرف Keras [۵۰] که به زبان پایتون نوشته شده است استفاده کردیم. این کتابخانه مجموعه‌ای مفید از توابع و API‌ها برای ساخت انواع مختلف شبکه‌های عصبی را در اختیار قرار می‌دهد؛ اما، برای اجرا نیاز به یک چارچوب سطح پایین دارد که ما Tensorflow [۴۷] را بدین منظور انتخاب کردیم. Tensorflow همچنین یک ابزار به نام Tensorboard دارد که از آن می‌توان برای مصورسازی گراف‌های محاسباتی و نیز مشاهده نمودارهای خطا و دقت در فرایندهای آموزش و آزمون استفاده کرد. برای آموزش مدل‌ها ما از تابع خطای CE (رابطه ۲-۶ در بخش ۲-۳-۳) و بهینه‌ساز Adam [۵۱] با نرخ‌های یادگیری 10^{-4} و 10^{-3} استفاده کردیم.

هدف این پایان‌نامه همان‌طور که پیش از این گفته شد، ارائه یک روش تولید خودکار داده آزمون است که در بخش‌های گذشته این فصل بحث شد. اما تولید داده آزمون به‌تنهایی کافی نیست و برای ارزیابی لازم است تا یک فازر قالب فایل با همه پیمانه‌های شکل ۲-۲ در اختیار داشته باشیم. برای این منظور نیاز به دو پیمانه تزریق کننده مورد آزمون و پایش SUT نیز داریم. اغلب نرم‌افزارهایی که یک فایل را به عنوان ورودی

الگوریتم ۲-۴ DataNeuralFuzz

Input: Learnt model M , Sequence prefix P , Diversity D , Fuzzing rate FR , End token ET , Binary token BT

Output: Test data TD

```

1  $TD \leftarrow P$ 
2  $MaxLen \leftarrow RandInt(a, b)$ 
3 while not EndsWith( $TD, ET$ ) do
4    $predicts \leftarrow Predict(M(P))$ 
5    $c, p(c) \leftarrow Sample(predicts, D)$  /* Sample  $c$  from the learnt model */
6    $p_{fuzz} \leftarrow Random(0, 1)$  /* Decide whether to fuzz */
7   if  $p_{fuzz} < FR \wedge p(c) < \alpha \wedge c \notin Chars(BT) \wedge c \notin Chars(ET)$  then
8      $c \leftarrow argmin_{c'} \{p(c') \in predicts\}$  /* Fuzz  $c$  by  $c'$  where  $c'$  is the
       lowest likelihood */
9   end
10   $TD \leftarrow TD + c$ 
11   $P \leftarrow P[1:] + c$  /* Propagate fuzz to prefix and next generated data */
12  if  $Len(TD) > MaxLen$  then
13     $TD \leftarrow TD + ET$ 
14    Break
15  end
16 end
17 if  $BT \in TD$  then
18    $TD \leftarrow AddBinaryPart(TD)$ 
19    $TD \leftarrow MutateBinaryPart(TD)$ 
20 end
21 Return  $TD$ 

```

الگوریتم ۳-۴ MetadataNeuralFuzz

Input: Learnt model M , Sequence prefix P , Diversity D , Fuzzing rate FR , End token ET , Binary token BT

Output: Test data TD

```

1  $TD \leftarrow P$ 
2  $MaxLen \leftarrow \text{RandInt}(a, b)$ 
3 while not EndsWith ( $TD$ ,  $ET$ ) do
4    $predicts \leftarrow \text{Predict}(M(P))$ 
5    $c, p(c) \leftarrow \text{Sample}(predicts, D)$       /* Sample  $c$  from the learnt model */
6    $p\_fuzz \leftarrow \text{Random}(0, 1)$           /* Decide whether to fuzz */
7   if  $p\_fuzz < FR \wedge p(c) > \beta$  then
8      $c' \leftarrow \text{argmin}_{c''} \{p(c'') \in predicts\}$  /* Fuzz  $c$  by  $c'$  where  $c'$  is the
9       lowest likelihood */
9   end
10   $TD \leftarrow TD + c'$ 
11   $P \leftarrow P[1:] + c$                     /* Don't propagate fuzz to prefix */
12  if  $Len(TD) > MaxLen$  then
13     $TD \leftarrow TD + ET$ 
14    Break
15  end
16 end
17 if  $BT \in TD$  then
18    $TD \leftarrow \text{AddBinaryPart}(TD)$ 
19    $TD \leftarrow \text{MutateBinaryPart}(TD)$ 
20 end
21 Return  $TD$ 

```


می‌پذیرند یک واسط خط فرمان نیز دارند که می‌توان از این طریق یک فایل را به آنها داد. در غیر اینصورت نیز می‌توان یک برنامه کوچک برای تزریق ورودی نوشت. به‌طور کلی در آزمون فازی قالب فایل تزریق ورودی ساده است. برای پایش نیز ابزارهای مستقلی وجود دارد که می‌توان از آنها استفاده کرد. ما آزمون فازی را تحت سیستم عامل ویندوز انجام دادیم^۱ و ابزار Application Verifier [۳۶] که در بخش ۲-۲-۲ معرفی شد را برای پایش انتخاب کردیم که مایکروسافت استفاده از آن را در SDL توصیه کرده است.

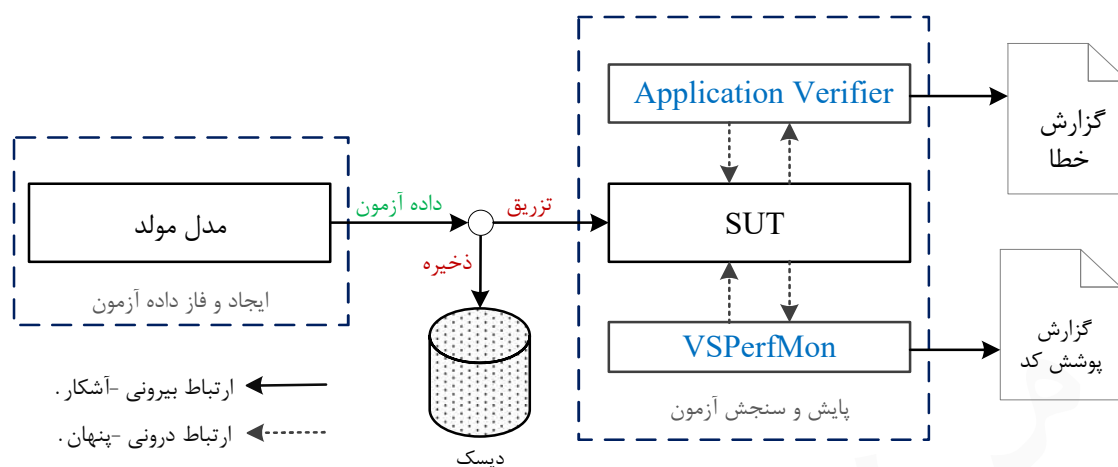
Application Verifier فایل اجرایی SUT را می‌گیرد و در هر بار اجرای SUT یک فایل حاوی وقایع را با یک شماره ترتیبی ثبت می‌کند. در صورتی که برنامه دچار خطای زمان اجرا (شامل یکی از انواع خطاهای فساد حافظه، دسترسی غیر مجاز و غیره) شود، نوع خطا در این فایل ثبت و ضبط می‌گردد. سپس می‌توان برنامه را به‌همراه داده آزمون مسبب خطا (که از روی شماره ترتیبی فایل وقایع قابل تشخیص است) در یک محیط اشکال‌زدا مانند WinDbg مجدداً اجرا و ضمن مشاهده خطا، محل دقیق آن را آشکار ساخت.

از سازوکاری که در بالا توضیح داده شد برای آزمون فازی نرم‌افزار MuPDF که یک نرم‌افزار پر استفاده است و قالب پیچیده PDF را به‌عنوان ورودی می‌پذیرد، در آزمایش‌ها استفاده کرده‌ایم. ما همچنین پوشش کد این نرم‌افزار را با استفاده از ابزار VSPerfMon اندازه‌گیری کردیم تا عملکرد تولید کننده داده آزمون را بسنجیم. شکل ۴-۵ یک طرح‌واره کلی از معماری (مؤلفه‌ها و ارتباط بین آنها) فازر پیشنهادی ما در این پایان‌نامه را نشان می‌دهد. با تعویض پیمانه پایش، این فازر در سیستم‌های عامل دیگر نیز قابل اجرا خواهد بود. این فازر را می‌توان در گروه فازرهای جعبه خاکستری، مبتنی بر روش ترکیبی تولید داده آزمون و بدون حلقه بازخورد به‌شمار آورد. VSPerfMon برای اندازه‌گیری پوشش کد در حالت جعبه سفید استفاده می‌شود که در صورت عدم وجود کد منبع SUT بایستی آن را با یکی از ابزارهای سنجش پوشش کد باینری در بخش ۲-۱-۳ جایگزین کرد.

۴-۵ خلاصه

مدل‌های زبانی عصبی در سطح کاراکتر قابلیت یادگیری یک توزیع آماری روی یک توالی از کاراکترها را دارند. در این فصل ما ساختار یک فایل را به صورت یک توالی از کاراکترها (بایت‌ها) مدل کرده و با ارائه چندین مدل زبانی عصبی و چندین راهبرد تولید داده‌های جدید از این مدل‌ها، یک روش تولید خودکار داده آزمون ارائه دادیم که قابل استفاده در فازرهای قالب فایل مختلف است. ما همچنین دو الگوریتم با عنوان‌های فاز داده عصبی و فاز فراداده عصبی ارائه کردیم. هر کدام از این الگوریتم‌ها بدشکل‌سازی داده آزمون ورودی

^۱ هرچند که روش پیشنهادی کاملاً مستقل از سیستم عامل است. برنامه تولید خودکار داده آزمون به زبان پایتون نوشته شده و روی سکوها‌های مختلف قابل اجرا است.



شکل ۴-۵: معماری فازر قالب فایل پیشنهادی. فازر به صورت کاملاً پیمانه‌ای پیاده‌سازی شده است. پیکان‌ها ارتباطات بین مؤلفه‌ها را نشان می‌دهند. ارتباط بیرونی، سازوکار ارتباطی است که توسط ما پیاده‌سازی شده یا قابل تغییر است و ارتباط درونی سازوکار ارتباطی است که توسط پیمانه‌های شخص ثالث استفاده می‌شود.

را با هدف خاصی انجام می‌دهند: اولی با هدف جابه‌جایی داده‌های یک فایل به منظور ایجاد خطا در هنگام پرداخت آنها و دومی با هدف جابه‌جایی فراداده‌های یک فایل برای اینجا خطا در تجزیه‌گر اولیه ساختار فایل تحت فاز.

در بخش پایانی فصل نیز یک فازر قالب فایل را طراحی و پیاده‌سازی کردیم که از آن می‌توان برای آزمون فازی برنامه‌هایی با ورودی فایل استفاده کرد. خروجی این فازر علاوه بر گزارش خطاهای یافت شده حین فرایند آزمون، گزارش میزان پوشش کد و ذخیره کلیه داده‌های آزمون تولید شده نیز است. معماری فازر پیشنهادی کاملاً پیمانه‌ای بوده و با تعویض پیمانه پایش، امکان استفاده از آن در سیستم‌های عاملی به جز ویندوز فراهم می‌شود. فازر معرفی شده یک فازر جعبه خاکستری و مجهز به یک روش تولید داده آزمون ترکیبی است. در فصل ۵، به ارزیابی روش پیشنهادی و بررسی یافته‌های حاصل از آن می‌پردازیم.

فصل ۵

ارزیابی روش پیشنهادی

«من شکست نخورده‌ام. من فقط ۱۰,۰۰۰ راه پیدا کردم که کار نمی‌کند.»

✠ توماس ادیسون

ما روش پیشنهادی خود را روی قالب فایل PDF به عنوان یک قالب فایل با ساختار پیچیده و ترکیبی و نرم افزار متن باز و رایگان MuPDF [۲۵] مورد ارزیابی قرار دادیم. در این فصل به تشریح چیدمان آزمایش ها، معیارهای ارزیابی، یافته های حاصل از آزمایش ها و نتیجه گیری از آنها می پردازیم. هدف علاوه بر اثبات کارایی و خوب بودن روش پیشنهادی، شناسایی و تفکیک پارامترهای مهم در هنگام استفاده از فنون یادگیری ماشینی در آزمون فازی و تولید خودکار داده آزمون است. استفاده از یادگیری ماشینی در آزمون فازی حوزه پژوهشی بدیع و ناشناخته ای است. بنابراین در این فصل معیارهایی را که در هر روش پیشنهادی بایستی مورد ارزیابی قرار بگیرند، نیز مشخص خواهیم نمود. سپس روش پیشنهادی خود را با آنها مورد ارزیابی قرار می دهیم.

۵-۱ مورد مطالعاتی

نرم افزار MuPDF [۲۵] مجموعه غنی از کتابخانه ها، API ها و ابزارهای کاربردی برای کار با فایل های PDF است. این نرم افزار مشتمل بر سه قسمت API، ابزار mutool و ابزار PDF خوان است. ابزار mutool برای ایجاد و تغییر فایل های PDF استفاده می شود. ابزار PDF خوان نیز برای نمایش فایل های PDF به کار می رود. بسته نرم افزاری کامل، به زبان C نوشته شده و در سیستم عامل های ویندوز، لینوکس و Mac و همچنین اندروید قابل استفاده است. به دلیل قابلیت حمل بالا و کاربرد زیادی که دارد، پیدا کردن خطا در آن حائز اهمیت است. در

آزمایش‌ها ما ابزار PDF خوان نسخه ویندوز را مورد آزمون قرار دادیم. منظور از عبارت MuPDF در ادامه این فصل همان ابزار PDF خوان است.

MuPDF البته در نسخه‌های جدید قالب‌های فایل دیگری را نیز پشتیبانی می‌کند؛ از جمله قالب‌های فایل XPS و EPub. اما قالب فایل PDF بدون شک رایج‌ترین قالب اسناد و انتشارات الکترونیکی است که بیشتر مردم از آن استفاده می‌کنند. از این رو تضمین کیفیت و امنیت PDF خوان‌ها به‌ویژه در تلفن‌های همراه اهمیت دو چندانی دارد. افزون بر آنچه گفته شد قالب فایل PDF و نرم‌افزار MuPDF تمامی مسائلی را که در بخش ۲-۱ بدان اشاره کردیم، یعنی ساختار پیچیده ورودی، ساختار پیچیده کد و غیره، دارند و از طرفی در فصل ۳ دیدیم که فازهای مبتنی بر جابه‌جایی مانند AFL و AFL افزوده نتوانسته بودند به پوشش کد خوبی حین آزمون فازی این نرم‌افزار دست‌یابند. بنابراین انتخاب آنها به‌عنوان مورد مطالعاتی به‌دلایل بالا و نیز در جهت فراهم شدن امکان مقایسه با کارهای پیشین، قابل توجه است. آزمایش‌ها برای هر قالب فایل و هر نرم‌افزار دیگری به‌نحوی که در ادامه این فصل توضیح داده می‌شود قابل اجرا است.

ساختار قالب فایل PDF در پیوست آ به‌طور کامل بررسی شده است. بخش عمده این ساختار را اشیای متنی تشکیل می‌دهند. از مدل‌های معرفی شده در بخش ۲-۴ برای یادگیری ساختار اشیای متنی استفاده می‌کنیم؛ زیرا یادگیری فایل PDF به‌صورت یکجا با توجه به اینکه قسمت‌هایی از آن برای آدرس‌دهی استفاده می‌شوند امکان‌پذیر نیست. اما MuPDF یا هر PDF خوان دیگری فایل‌های کامل PDF را به عنوان ورودی می‌پذیرد به همین لازم است تا اشیای تولید شده را به یک فایل تبدیل کنیم. برای این کار اشیای جدید را با استفاده از سازوکار توضیح داده شده در پیوست آ در مورد بروزرسانی یک فایل PDF، به یک فایل معتبر از پیش موجود، که آن را میزبان^۱ می‌نامیم، اضافه می‌کنیم. در نتیجه یک فایل PDF جدید در اختیار خواهیم داشت که ساختار آن معتبر بوده و برخی اشیای آن تغییر داده شده و بازنویسی شده‌اند.

۲-۵ معیارهای ارزیابی

در تولید خودکار داده آزمون به روش یادگیری ژرف، تعداد زیادی پارامتر وجود دارد که می‌توان تأثیر آنها را در فرایند آزمون فازی سنجید. اما مهم‌ترین هدف در فرایند آزمون فازی همان‌گونه که پیش از این هم ذکر شد یافتن خطا در SUT است. هدف مهم بعدی افزایش میزان پوشش کد برای نیل هر چه بیشتر به هدف اول است. هدف دیگر که مختص به این روش است یادگیری هر چه بهتر ساختار فایل ورودی است. براساس اهداف اشاره شده، معیارهای ارزیابی زیر را در این فصل لحاظ کرده‌ایم:

- **خطا و دقت مدل‌های یادگیری ژرف.** شامل خطا و دقت هر مدل روی مجموعه‌های آموزش و ارزیابی،

¹Host

میزان و نحوه تغییر آنها با گذشت زمان. این معیارها توسط کد نوشته شده به ابزار Tensorboard گزارش شده و در حین فرایند آموزش یا در پایان آن قابل مشاهده هستند.

● **سرگشتگی.** برای ارزیابی میزان خوب بودن مدل زبانی در پیش‌بینی نشانه بعدی توالی داده شده. معیار سرگشتگی از رابطه ۲-۱۴ و با جایگذاری خطای مدل روی مجموعه ارزیابی، برای آن مدل محاسبه می‌گردد.

● **پوشش کد.** در هر اجرای SUT، تعداد خطوط برنامه و نیز تعداد بلوک‌های اولیه‌ای که اجرا شده‌اند، شمارش می‌شود. با داشتن کل خطوط و بلوک‌های اولیه درصد پوشش کد نیز قابل محاسبه است. برای یک مجموعه آزمون، یعنی چندین بار اجرای متوالی SUT پوشش کد عبارت خواهد بود از اجتماع پوشش کد تک‌تک هریک از اجراها. بنابراین منظور از اصطلاح کلی پوشش کد در ادامه پوشش بلوک‌های پایه است.

● **خطا و آسیب‌پذیری.** SUT تحت ابزار پایش Application Verifier اجرا می‌شود که خطاهای احتمالی را ثبت و شماره‌گذاری می‌کند. سپس می‌توان با تحلیل این خطاها در یک محیط اشکال زدا، ضمن تعیین مکان و علت خطا، امکان قابل بهره‌برداری بودن آنها را بررسی کرد.

آزمایش‌ها بایستی به‌گونه‌ای طراحی شوند که هریک از آنها تأثیر یک پارامتر خاص را بر روی معیارهای بالا، به‌طور واضح نمایان سازد؛ اما، پارامترها کدامند؟ ما در این بخش، چندین پارامتر و مقادیر قابل آزمایش برای هریک را شناسایی و معرفی می‌کنیم. فهرستی از این پارامترها و مقادیری که می‌توانند بپذیرند در جدول ۵-۱ ذکر شده است. همان‌طور که مشاهده می‌شود برخی از پارامترها مقادیری پیوسته و برخی دیگر مقادیری گسسته می‌پذیرند. بر این باور هستیم که هر روش تولید خودکار داده آزمون با استفاده از یادگیری ماشینی روی چنین پارامترهایی بنا می‌شود، لذا شناسایی و تفکیک این پارامترها پیش از هرگونه طرح آزمایشی لازم و ارزشمند است.

نوع مدل در سطر اول جدول ۵-۱ برای هر مدل، خود دارای ابرپارامترهایی مانند تعداد عصب، لایه، نرخ یادگیری و غیره است که در فصل قبل به آنها اشاره شد. بنابراین جدول ۵-۱ یک دید جامع از همه پارامترهایی که تا اینجا معرفی شدند، در اختیار می‌گذارد. در بخش بعد چیدمان آزمایش‌ها خود را براساس این پارامترها شرح خواهیم داد.

جدول ۵-۱: پارامترهای مؤثر در تولید خودکار داده آزمون با روش‌های یادگیری ماشینی و مقادیر قابل آزمایش برای آنها.

ردیف	پارامتر	مقادیر قابل آزمایش
۱	نوع مدل	• انواع مدل‌های یادگیری ژرف
۲	راهبرد	• حریصانه • نمونه برداری • ترکیبی (حریصانه + نمونه برداری) • چند فهرست بهتر
۳	تنوع (D)	• $(0, \infty)$
۴	روش فاز	• بدون فاز • تصادفی • مبتنی بر مدل • ترکیبی
۵	الگوریتم فاز	• فاز داده • فاز فراداده
۶	تعداد دوره آموزش	• $\{1, 2, 3, \dots\}$
۷	نوع فایل	• متنی • متنی + دودویی
۸	استفاده از میزبان	• استفاده • عدم استفاده
۹	تعداد میزبان	• $\{1, 2, 3, \dots\}$
۱۰	میزان تغییر میزبان	• ثابت • متغیر

جدول ۵-۲: مشخصه‌ها و زمان آموزش مدل‌های جدول ۴-۱ در فرایند آموزش.

پارامتر / شماره مدل	۱	۲	۳	۴
طول توالی‌های ورودی (d)	۵۰	۵۰	۵۰	۵۰
گام پرش (j)	۳	۳	۱	۱
تعداد دوره	۵۰	۵۰	۵۰	۵۰
زمان تقریبی یک دوره (ساعت: دقیقه)	۱:۰۰'	۱:۴۵'	۵:۴۵'	۹:۳۰'
حجم تقریبی مدل (مگابایت)	۱/۲۴	۲/۷۶	۹/۹۹	۵/۴۱

۳-۵ چیدمان آزمایش‌ها

برای آموزش مدل‌های یادگیری معرفی شده در بخش ۴-۲ (جدول ۴-۱) و تولید داده‌های آزمون از یک سیستم فیزیکی با پردازنده گرافیکی NVidia GTX 1080، پردازنده مرکزی Intel Core i7 و ۲۰ گیگابایت حافظه اصلی به همراه سیستم عامل Ubuntu 16.04 x64 استفاده کردیم. آزمون‌های فازی را نیز روی ماشین مجازی با پردازنده Intel Core i7 و ۸ گیگابایت حافظه اصلی به همراه سیستم عامل Windows 10 x64 انجام دادیم. ما همچنین از نسخه نهایی MuPDF در زمان انجام آزمایش‌ها یعنی نسخه (2017-04-11) MuPDF 1.11 Final برای آزمون فازی استفاده کردیم. باتوجه به نهایی بودن نسخه انتظار می‌رود خطاهای نسخه تا حد زیادی، در مقایسه با نسخه‌های آلفا، بتا و RC برطرف شده باشند و لذا شناسایی خطا کار سخت‌تری خواهد بود. مشخصه‌های کامل مدل‌ها در فرایند آموزش به همراه تعداد دوره و زمان آموزش هر دوره، در جدول ۵-۲ آمده است.

پیچیدگی مدل‌های جدول ۵-۲ با افزایش شماره، افزایش می‌یابد. یعنی مدل ۲ از مدل ۱ پیچیده‌تر در نظر گرفته شده است. همان‌طور که در جدول هم مشخص است مدل‌های پیچیده‌تر زمان آموزش طولانی‌تری دارند. برای مدل‌های پیچیده‌تر منطقی است که نمونه‌های آموزشی بیشتری داشته باشیم. در نتیجه برای مدل‌های ۳ و ۴ گام پرش را ۱ و برای مدل‌های ۱ و ۲ گام پرش را ۳ در نظر گرفتیم. همچنین برای مدل شماره ۳ از روش منظم‌سازی Dropout [۴۲] $p = 0.3$ استفاده کردیم. در نهایت هر مدل را دست‌کم برای ۵۰ دوره روی داده‌های مجموعه آموزش، آموزش دادیم.

۵-۳-۱ مجموعه داده

آموزش موفق شبکه‌های عصبی ژرف مستلزم داشتن مجموعه داده به اندازه کافی بزرگ و مناسب است. در یادگیری ماشینی هرچه قدر تعداد داده‌ها بیشتر باشد، یادگیری بهتر انجام خواهد شد. هنوز در بسیاری از زمینه‌ها

مجموعه داده کافی برای آموزش وجود ندارد و این یکی از محدودیت‌های استفاده از شبکه‌های عصبی ژرف است. برای یادگیری آماری ساختار اشیای فایل PDF پیکره بزرگی از فایل‌های PDF جمع‌آوری کردیم؛ زیرا چنین پیکره‌ای از قبل وجود نداشت. بخشی از این پیکره شامل مجموعه داده‌های آزمون PDF خوان Mozilla^۱ است که در مرورگر وب Firefox و دیگر پروژه‌های این شرکت استفاده می‌شود و به صورت رایگان در دسترس است. بخش دیگری از آن PDFهای استفاده شده در فازهای دیگر مثل AFL است و بلاخره بخشی بزرگی نیز از طریق وب جمع‌آوری شد به نحوی که تنوع خوبی از لحاظ اندازه و محتوا داشته باشند. پیکره نهایی در حال حاضر شامل ۶۱۰۰ فایل PDF مختلف و طبقه‌بندی شده است که ما آن را تحت نام IUST PDF Corpus^۲ منتشر کرده‌ایم و از آن می‌توان در موارد دیگر نیز استفاده کرد.

از فایل‌های داخل این پیکره که اندازه آنها بین ۱ تا ۹۰۰ کیلوبایت و تا حداکثر ۷۹۳۵ کیلوبایت متغیر است، در حدود ۵۰۴۱۵۳ تعداد شیء داده‌ای PDF با طول‌های گوناگون استخراج شد که به عنوان مجموعه داده جهت یادگیری ساختار اشیای داده‌ای PDF در نظر گرفته شده‌اند. این اشیاء سپس همگی در یک فایل الحاق شدند که حجم فایل نهایی در حدود ۷۰ مگابایت است. از این تعداد در حدود ۱۳۷۱۵۷ شیء حاوی جریان‌های دودویی بودند که محتوای همه این جریان‌ها پس از شناسایی با توکن دودویی stream جایگزین شدند؛ زیرا، همان‌طور که قبل از این هم اشاره کردیم نیازی به شرکت دادن آنها در فرایند آموزش نیست و بعداً مجدداً آنها را اضافه می‌کنیم.

۵-۳-۲ پیش‌پردازش داده‌ها

تعداد اشیای استخراج شده بسیار زیاد هستند. قبل از انجام عملیات یادگیری مدل بایستی تعدادی از این اشیاء را به نحوی از مجموعه داده حذف کرد. برای این منظور عملیات پیش‌پردازی طی مراحل مختلف روی مجموعه داده انجام دادیم. مهمترین ویژگی در دسترس در تنظیم ابرپارامترهای مدل یادگیری ژرف در اینجا طول اشیای داده‌ای PDF است که هم در فرایند آموزش و هم در فرایند تولید نقش تأثیرگذاری دارد. لذا داشتن توزیعی همگن از اشیای PDF برحسب طول آنها فرضیه‌ای است که در اینجا در نظر گرفته شده است و در عملیات پیش‌پردازش نیز معیار طول اشیای داده‌ای بوده است. ویژگی طول اشیاء از این منظر مهم است که شبکه LSTM بر روی توالی‌های بسیار طولانی عملکرد ضعیفی دارد و بهتر است چنین توالی‌هایی را در نظر نگیریم. مراحل پیش‌پردازش داده‌ها به ترتیب در زیر ذکر شده است.

۱. هر شیء داده‌ای یک موجودیت مستقل در نظر گرفته شده و ابتدا مجموعه داده برحسب طول اشیای

^۱<https://github.com/mozilla/pdf.js/tree/master/test/pdfs>

^۲https://github.com/m-zakeri/iust_deep_fuzz/tree/master/dataset

داخل آن مرتب گردید. سپس کلیه اشیاء داده‌ای به شکل

```
obj
/
endobj
```

و

```
obj
null
endobj
```

از مجموعه داده اولیه حذف شدند؛ زیرا، حاوی بار اطلاعاتی مفیدی نیستند.

۲. پس از اعمال مرحله ۱ تعداد اشیای مجموعه داده به ۴۹۴۹۷۹ کاهش یافت. اکنون به طور تقریبی هر صدک از مجموعه داده شامل حدود ۴۹۵۰ شیء داده‌ای PDF است. برای حذف هرچه بیشتر داده‌های پرت صدک‌های اول و آخر را نیز از مجموعه داده حذف کردیم. به این ترتیب در پایان این مرحله تعداد ۴۸۵۰۸۰ شیء داده‌ای PDF باقی ماند.

۳. مرحله اصلی پیش پردازش خارج کردن داده‌های outlier و extreme value با استفاده از معیار *iqr* توسط ابزار WEKA^۱ [۵۲] است. ورودی این مرحله مجموعه داده مرحله ۲ است و خروجی آن حاوی تعداد ۴۷۷۱۰۴ شیء داده‌ای PDF است. برای تعیین دو مقدار نامبرده یعنی outlier و extreme value به ترتیب از ضرایب ۳ برای ضریب outlier و ۶ برای ضریب extreme value استفاده شد؛ که مقادیر پیش فرض WEKA برای پاک‌ساز *iqr*^۲ هستند.

۴. در این مرحله هدف تقسیم مجموعه داده به سه بخش مجموعه آموزش، مجموعه ارزیابی و مجموعه آزمون است. از آنجایی که یادگیری مدل ما در حالت بدون نظارت انجام می‌شود^۳، از مجموعه آزمون فقط برای انتخاب پیشوند جهت تولید اشیای جدید استفاده خواهد شد. به عبارت دیگر مدل روی مجموعه آموزش، آموزش داده می‌شود و در هنگام تولید اشیای جدید PDF مقادیر اولیه از مجموعه آزمون برداشته می‌شود. برای این منظور ۲۵ درصد از کل تعداد اشیای PDF را برای مجموعه آزمون و ۷۵ درصد را برای مجموعه آموزش/ارزیابی در نظر گرفتیم. برای ایجاد این تفکیک نیز از فیلتر

^۱ <https://www.cs.waikato.ac.nz/ml/weka/>

^۲ Filter

^۳ برچسب‌ها از خود مجموعه آموزش برداشته می‌شوند.

جدول ۵-۳: نحوه و درصدهای تقسیم مجموعه داده به مجموعه‌های آموزش، ارزیابی و آزمون.

مجموعه	تعداد شیء	درصد
آموزش	۲۶۸۳۷۱	۵۶/۲۵
ارزیابی	۸۹۴۵۷	۱۸/۷۵
آزمون	۱۱۹۲۷۶	۲۵
جمع	۴۷۷۱۰۴	۱۰۰

StratifiedRemoveFolds موجود در ابزار WEKA با دانه اولیه ۱۰۱ (کوچک‌ترین عدد اول سه رقمی) و مقدار $fold = ۳$ (یعنی بعد از درهم‌آمیزی^۱ و تقسیم کل مجموعه داده به ۴ قسمت با نسبت‌هایی برابر با مجموعه داده اولیه، قسمت سوم آن برای مجموعه آزمون کنار گذاشته شود) استفاده شد. در نهایت تعداد ۳۵۷۸۲۸ شی در مجموعه‌های آموزش/ارزیابی و تعداد ۱۱۹۲۷۶ شی داده‌ای هم در مجموعه آزمون قرار گرفتند.

۵. در تکمیل مرحله ۴، نیاز است تا مجموعه آموزش/ارزیابی جدا شده، حاوی ۳۵۷۸۲۸ شیء خود به دو مجموعه مجزای آموزش و ارزیابی تقسیم شود. برای این منظور این مجموعه آموزش/ارزیابی به ۴ قسمت تقسیم گردید و قسمت چهارم به عنوان مجموعه ارزیابی در نظر گرفته شد. سه قسمت اول هم مجموعه آموزش را تشکیل می‌دهند. درصد تعداد اشیاء هریک از مجموعه‌های آموزش، ارزیابی و آزمون در جدول ۵-۳ آمده است.

۶. آخرین مرحله پیش‌پردازش نرمال‌سازی^۲ متن است. در این مرحله برخی کاراکترهای بسیار کم تکرار (فراوانی نسبی نزدیک صفر) حذف یا با کاراکترهای مناسبی جایگزین شدند. در واقع طی این مرحله اندازه مجموعه واژگان کاراکترها از ۹۶ به ۶۴ تقلیل یافت. علت انجام این مرحله، آن است که کاراکترهای با تعداد کم داده پرت محسوب می‌شوند و به دلیل اینکه در مجموعه آموزش تعداد تکرار آنها پایین است عملاً احتمال بسیار ناچیزی به آنها تخصیص داده می‌شود که در واقع ضرورتی به این کار نیست. در عین حال با انجام این کار، حجم مکعب داده‌ها کمتر شده و فرایند آموزش سریع‌تر انجام می‌شود.

^۱Shuffle

^۲Normalization

۵-۳-۳ انتخاب فایل‌های میزبان

در [۱۱] انتخاب فایل‌های میزبان به صورت انتخاب سه فایل با کمترین حجم از پیکره انجام گرفته است. در واقع این انتخاب بدون هدف خاص و کاملاً تصادفی صورت گرفته است. پیرو روش آنها ما نیز برای آزمایش‌های خود سه فایل را به عنوان فایل‌های میزبان انتخاب کردیم، با این تفاوت که برای انتخاب فایل‌های میزبان ابتدا میزان پوشش کد را به صورت مجزا برای همه فایل‌های IUST PDF Corpus به دست آورده و سپس سه فایل با بیشترین میزان پوشش کد، کمترین میزان پوشش کد و نیز میزان پوشش کد میانگین، به ترتیب با نام‌های $host1_max$ ، $host2_min$ و $host3_avg$ به عنوان فایل‌های میزبان انتخاب شدند. هدف از این کار بررسی تأثیر روش پیشنهادی بر بهبود میزان پوشش کد در فایل‌های با اختلاف پوشش کد زیاد است و سپس انجام آزمون فازی براساس بهترین میزبان است. ما همچنین برای هر آزمایش پوشش کد مجموع هر سه میزبان را با ادغام پوشش کدهای هریک از آنها تحت عنوان $host123$ گزارش کردیم. این کار کمک می‌کند تا بتوانیم نتایج دو آزمایش کاملاً جدا را در حالت کلی با یدکدیگر مقایسه کنیم.

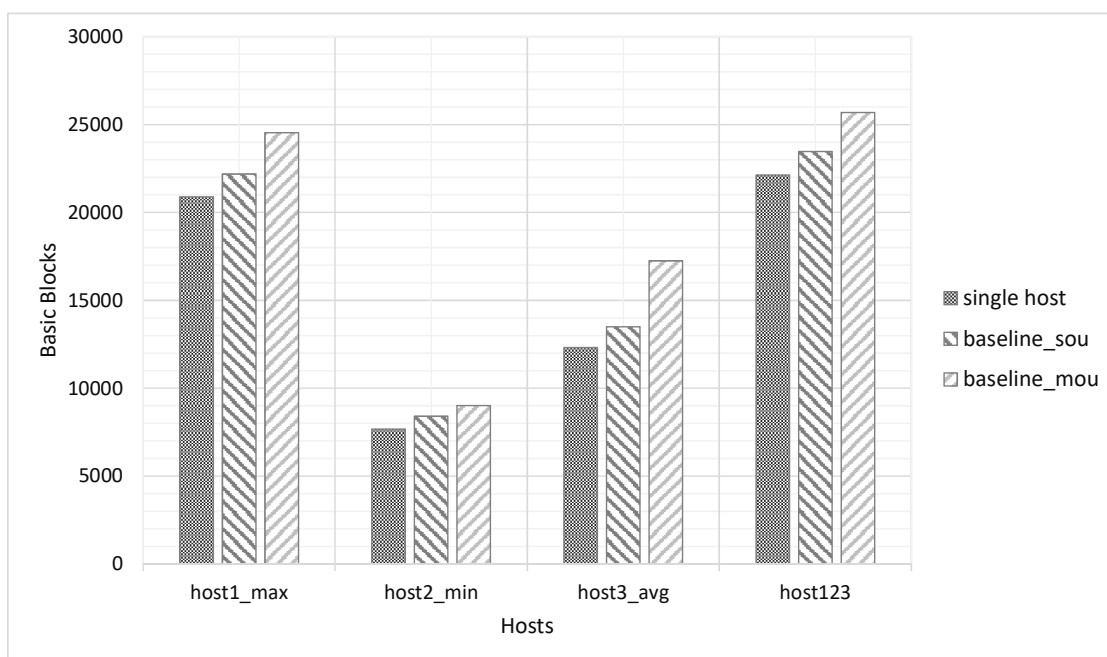
۵-۳-۴ پوشش کد مبنا

برای آنکه بتوانیم درک معنادار و مقایسه‌پذیری از پوشش کدهای به دست آمده در هر آزمایش داشته باشیم. ابتدا پوشش کد مبنا را مطرح، اندازه‌گیری و گزارش می‌کنیم. پوشش کد مبنا، پوشش کد در غیاب استفاده از روش پیشنهادی برای تولید خودکار داده آزمون است. برای این منظور، ابتدا از مجموعه آزمون، تعداد ۱۰۰۰ شیء داده‌ای PDF استخراج و در دو حالت به فایل‌های میزبان تزریق می‌نماییم:

- **بروزرسانی یک شیء (SOU^۱):** برای هر فایل میزبان در هر مرحله یک شیء را انتخاب و آخرین شیء میزبان را با آن بازنویسی می‌کنیم، یعنی در واقع تغییر تنها یک شیء از فایل میزبان در هر مرحله بدون توجه به اندازه و تعداد اشیای فایل میزبان.

- **بروزرسانی چندین شیء (MOU^۲):** تغییر درصد ثابتی از اشیای فایل میزبان در هر مرحله. در این روش ما برای $host1_max$ تعداد $\frac{1}{8}$ ، برای $host2_min$ تعداد $\frac{1}{4}$ و برای $host3_avg$ تعداد $\frac{1}{4}$ اشیای هر فایل را به صورت تصادفی انتخاب و با اشیای انتخابی از مجموعه آزمون بازنویسی کردیم. یعنی در این روش هر میزبان را به نسبت بیشتری تغییر می‌دهیم، به امید اینکه این تغییرات سبب افزایش پوشش کد شود. می‌توان اشیاء را باروش‌هایی غیر از روش تصادفی نیز برای بازنویسی انتخاب کرد؛ به عنوان مثال

^۱ Single Object Update^۲ Multiple Objects Update



شکل ۵-۱: پوشش کد برای هریک از میزبان‌ها و پوشش کدهای مبنا.

انتخاب براساس ترتیب نزولی یا صعودی ID هر شیء یا براساس بزرگی آدرس آنها. در این پایان‌نامه ما فقط روش تصادفی را آزمایش کردیم.

به این ترتیب سه مجموعه ۱۰۰۰ فایل از حالت اول و سه مجموعه ۱۰۰۰ فایل از حالت دوم حاصل می‌گردد که آنها را به ترتیب ^۱baseline_sou و ^۲baseline_mou می‌نامیم. پوشش کد هر مجموعه و نیز پوشش کد اجتماع هر سه میزبان به عنوان یک مقدار پایه برای آزمایش‌های بعدی اندازه‌گیری و تعیین شدند. نمودار شکل ۵-۱ مقادیر پوشش کد مبنا را در مقایسه با پوشش کد هریک از میزبان‌ها به صورت تنها نشان می‌دهد. لازم به ذکر است که هر یک از میزبان‌ها پیش از بروزرسانی‌های افزایشی برای تولید داده جدید بررسی شدند و معتبر بودن ساختار آنها تأیید گردید. لذا، صرف انجام این کار به منزله آزمون فازی نخواهد بود؛ زیرا، هنوز داده‌ها را بدشکل (فاز) نکرده‌ایم. هرچند ممکن است در این فرایند نیز خطایی قابل شناسایی باشد.

مشاهدات

- میزان پوشش کد هریک از baseline‌ها از پوشش کد فایل میزبان به تنهایی بیشتر است. یعنی تغییر میزبان‌ها منجر به اجرای بلوک‌های پایه جدید و احتمالاً مسیرهای اجرایی جدیدی، شده است.

^۱Baseline for Single Object Update

^۲Baseline for Multiple Object Update

- میزان پوشش baselineها متناسب با میزان‌ها است. یعنی به‌عنوان مثال host1_max بیشترین پوشش کد را در هر سه حالت تنها، baseline_sou و baseline_mou داشته است. این بدین معنی است که انتخاب فایل میزان بسیار حائز اهمیت است و بخش قابل توجهی از پوشش کد در هر حالت مربوط به فایل میزان است.
- میزان پوشش کد baseline_mou در تمامی میزان‌ها از baseline_sou بیشتر است. این بدین معنی است که با افزایش میزان تغییر فایل شاهد افزایش میزان پوشش کد بوده‌ایم.
- بیشترین پوشش کد مربوط به اجتماع پوشش کدهای baseline_mou در host123 است که نشان می‌دهد هر کدام از میزان‌ها مجموعه دستورات متفاوتی را اجرا کرده‌اند.
- در نهایت اعداد و ارقام پوشش کد در محدوده ۲۵ هزار بلوک پایه بوده که نشان می‌دهد MuPDF نرم‌افزاری پیچیده است.

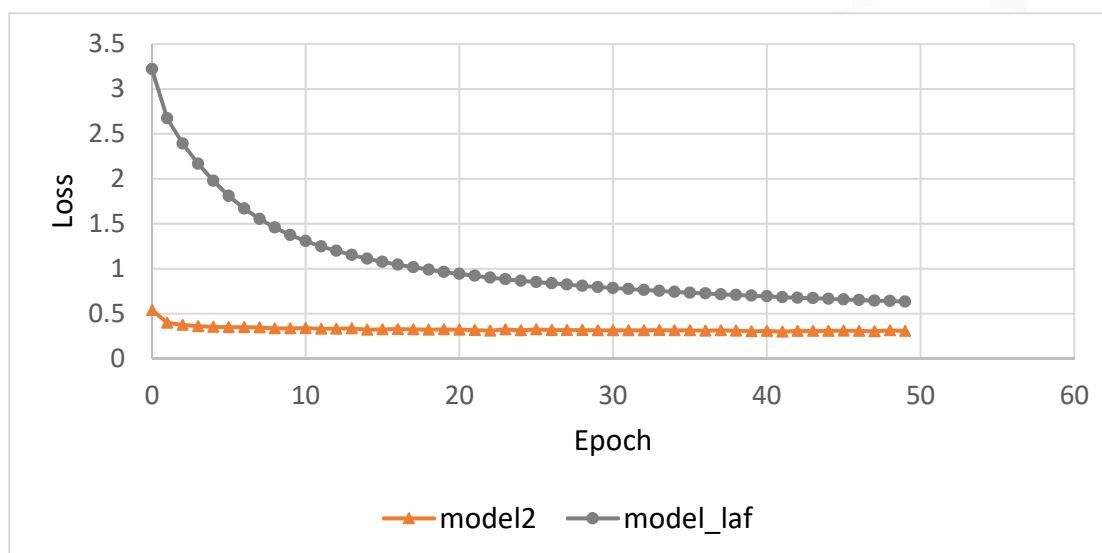
۴-۵ آزمایش‌ها، یافته‌ها و مقایسه نتایج

در این بخش آزمایش‌های صورت گرفته و نتایج حاصل از آن را بیان می‌کنیم. همچنین نتایج را با سایر کارهای مرتبط مقایسه خواهیم کرد. چون SUT مورد استفاده در یادگیری و فاز [۱۱] برای آزمون در دسترس نبود، ما روش یادگیری و فاز را نیز پیاده‌سازی و آن را روی ابزار MuPDF آزمایش کردیم. به این ترتیب توانستیم مقایسه معناداری بین روش پیشنهادی خود با این روش، به‌عنوان مرتبط‌ترین کار انجام شده داشته باشیم. نتایج نشان می‌دهد که روش پیشنهادی ما در معیار پوشش کد و همچنین در معیارهای خطا، دقت و سرگشتگی بهبود داشته است و توانسته روش یادگیری و فاز، که در آزمایش‌ها با نام laf آن را نشان می‌دهیم، را شکست دهد. روش پیشنهادی ما همچنین نسبت به AFL و AFLافزوده بهبودهایی داشته است که خواهیم دید.

چون مجموعه داده در مدل‌های یادگیری ژرف بسیار مهم است و روی نتایج تأثیر می‌گذارد در هنگام آموزش مدل laf مجموعه داده را مطابق مقیاس ارائه شده در [۱۱] لحاظ کردیم. بدین زمان آموزش کمتر برای هر دوره در حدود ۱۵ دقیقه به طول انجامید. سایر پارامترها مانند طول توالی‌های آموزشی d را که در مقاله به آن اشاره نشده بود نیز مشابه مقادیر آن در روش پیشنهادی قرار دادیم.

جدول ۴-۵: سرگشتگی، دقت و خطای مدل‌های مختلف روی مجموعه‌های آموزش و ارزیابی.

معیار / مدل	۱	۲	۳	۴	laf
سرگشتگی	۱/۴۴۰	۱/۳۹۱	۱/۳۳۵	۱/۳۵۰	۱/۸۶۰
بیشینه دقت آموزش	۰/۸۸۶	۰/۹۰۲	۰/۸۹۳	۰/۹۰۹	۰/۸۲۰
بیشینه دقت ارزیابی	۰/۸۸۴	۰/۸۹۵	۰/۹۰۴	۰/۹۰۵	۰/۸۰۰
کمینه خطای آموزش	۰/۳۵۳	۰/۲۹۸	۰/۳۲۴	۰/۲۷۶	۰/۶۲۳
کمینه خطای ارزیابی	۰/۳۶۵	۰/۳۳۰	۰/۲۸۹	۰/۳۰۰	۰/۷۲۴



شکل ۵-۲: نمودار تغییرات خطای مدل‌های ۲ و laf در دوره‌های ۱ تا ۵۰.

۴-۵-۱ سرگشتگی، خطا و دقت مدل‌ها

جدول ۴-۵ سرگشتگی، خطا و دقت مدل‌های چهارگانه پیشنهادی و مدل یادگیری و فاز را بعد از گذشت ۵۰ دوره، روی مجموعه‌های آموزش و ارزیابی نشان می‌دهد. سرگشتگی مطابق رابطه ۲-۱۴ و خطا نیز از رابطه ۲-۶ برای مدل‌ها محاسبه شده است که در فصل ۲ آنها را بیان کردیم. دقت نیز معیاری است که توسط Keras برای هر مدل حساب می‌شود. همچنین شکل ۵-۲ روند تغییرات خطا در فرایند آموزش را برای مدل ۲ و مدل laf نشان می‌دهد. مدل ۲ در این نمودار به این دلیل انتخاب گردیده که بیشترین میزان شباهت را از لحاظ معماری و مقادیر آبرپارامترها با مدل laf دارد.

مشاهدات

● سرگشتگی و خطای همه مدل‌های پیشنهادی از مدل یادگیری و فاز کمتر و دقت آنها از یادگیری و فاز بیشتر است. یعنی مدل زبانی عصبی با معماری توضیح داده شده در بخش ۴-۲، در یادگیری ساختار فایل از مدل کدگذار-کدگشا [۱۱] بهتر عمل کرده است.

● بیشترین دقت در بین تمامی مدل‌ها مربوط به مدل ۴ است. مدل ۴ LSTM دوسویه است. شبکه LSTM دوسویه هنگام آموزش علاوه بر کاراکترهای قبلی، کاراکترهای بعدی را نیز در نظر می‌گیرد. بنابراین توانسته است به دقت بیشتر دست یابد. کمترین میزان سرگشتگی نیز متعلق به مدل ۳ است. این مدل همان‌طور که در جدول نیز مشخص است دارای کمترین خطای ارزیابی است. البته اختلاف بین مقادیر مختلف در همه مدل‌های پیشنهادی کم است که نشان می‌دهد همه مدل‌ها قادر به یادگیری و درک ساختار فایل بوده‌اند. سرگشتگی بیشینه، حالت بدون استفاده از مدل یادگیری، برای تنظیمات ما برابر ۶۴ (اندازه بردار واژگان) خواهد بود.

● در شکل ۵-۲، منحنی خطای مدل ۲ در همه دوره‌ها، از منحنی خطای مدل laf پایین‌تر است. البته دوره‌ها زمان متفاوتی دارند بنابراین مقایسه نظیربه‌نظیر آنها ممکن است جالب به نظر نیاید. با این حال در بازه زمانی مساوی از شروع فرایند آموزش نیز این وضعیت فوق‌برقرار است. برای مثال پایان دوره ۱ برای مدل ۲ مصادف با پایان دوره ۱۲ برای مدل laf خواهد بود که باز هم خطای مدل laf بیشتر است. ما فرایند آموزش را برای هر دو مدل تا دوره ۱۰۰ ادامه دادیم و تغییر خلاقی مشاهده نشد. چنان‌چه در بخش ۴-۴-۵ خواهیم دید پوشش کد مدل ۲ نیز در دوره‌های مختلف از مدل laf بالاتر است.

۴-۵-۲ پوشش کد مدل‌های مولد

در آزمایش این بخش، مدل‌های مولد پیشنهادی خود را برای اولین مرحله تولید داده آزمون محک می‌زنیم. برای تولید داده‌های آزمون جدید از روش نمونه‌برداری و پارامتر تنوع D در همه آزمایش‌ها استفاده کردیم. همچنین هربار با یک پیشوند تصادفی از مجموعه آزمون تولید داده از مدل را شروع می‌کنیم. می‌خواهیم ببینیم برای میزبان‌های مختلف کدام مدل مولد و نیز کدام تنوع تولید داده، به پوشش کد بالاتری دست می‌یابد. هدف مقایسه کارکرد مدل‌ها و انتخاب مناسب‌ترین مدل برای استفاده در آزمون فازی است. برای این منظور هریک از مدل‌های جدول ۴-۱ را به تعداد ۵۰ دوره آموزش دادیم. در پایان هر دوره یک نمونه از مدل آموزش دیده را ذخیره کردیم. سپس مدل با کمترین خطا را از بین نمونه‌های ذخیره شده، برای تولید داده برگزیدیم.

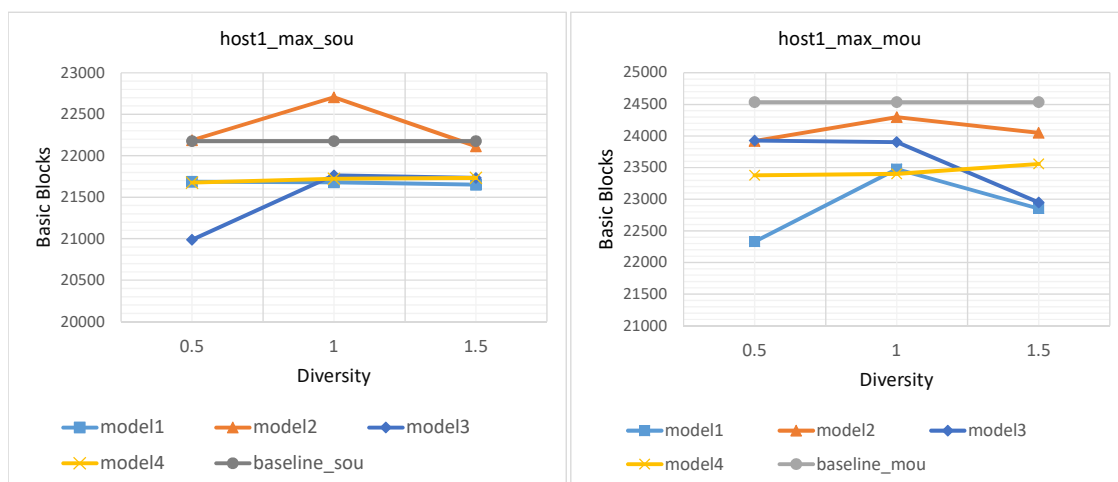
در مرحله بعد برای هر مدل در هنگام تولید داده با هریک از تنوع‌های تولید ۱،۰/۵ و ۱/۵، تعداد ۱۰۰۰ فایل PDF با هر میزبان تولید کرده‌ایم. یعنی در مجموع با هر مدل و هر میزبان ۳۰۰۰ فایل را تولید می‌کنیم. چون ۴ مدل و ۳ میزبان داریم، در کل ۳۶۰۰۰ فایل تولید می‌شود و چون دو روش بروزرسانی شیء SOU و MOU هم نیاز است در نهایت ۷۲۰۰۰ فایل تولید می‌کنیم. دقت شود که در این مرحله هنوز عملیات آزمون فازی یعنی فاز ورودی، مشاهده رفتار برنامه در اشکال‌زدا و ذخیره خطاهای احتمالی، اعمال نشده است و تنها تأثیر مدل‌ها، تنوع تولید داده از آنها و میزبان‌ها، بر روی میزان پوشش کد مورد نظر است.

برای مقایسه با baseline_sou تعداد ۱۰۰۰ شیء داده‌ای از هر مدل تولید کرده و برای مقایسه با baseline_mou تعداد ۳۰۰۰ شیء داده‌ای با هر مدل تولید می‌کنیم. تولید ۱۰۰۰ فایل PDF جدید با استفاده از مدل‌ها برای حالت SOU (با اندازه بافر ۱۰۰^۱) به‌طور میانگین ۶۰ دقیقه و برای حالت MOU به‌طور میانگین ۱۹۰ دقیقه به‌طول انجامید. در پایان پوشش کد نرم‌افزار MuPDF روی هریک از ۷۲ مجموعه ۱۰۰۰ فایل اندازه‌گیری شد. ما همچنین ادغام پوشش‌های کد را برای هر تنوع در قالب host123 اندازه گرفتیم. اندازه‌گیری پوشش کد برای هر مجموعه نیز در حدود ۶۰ دقیقه زمان برد. کلیه نتایج به تفکیک میزبان‌ها در شکل‌های ۳-۵ تا ۳-۵ ذکر شده است.

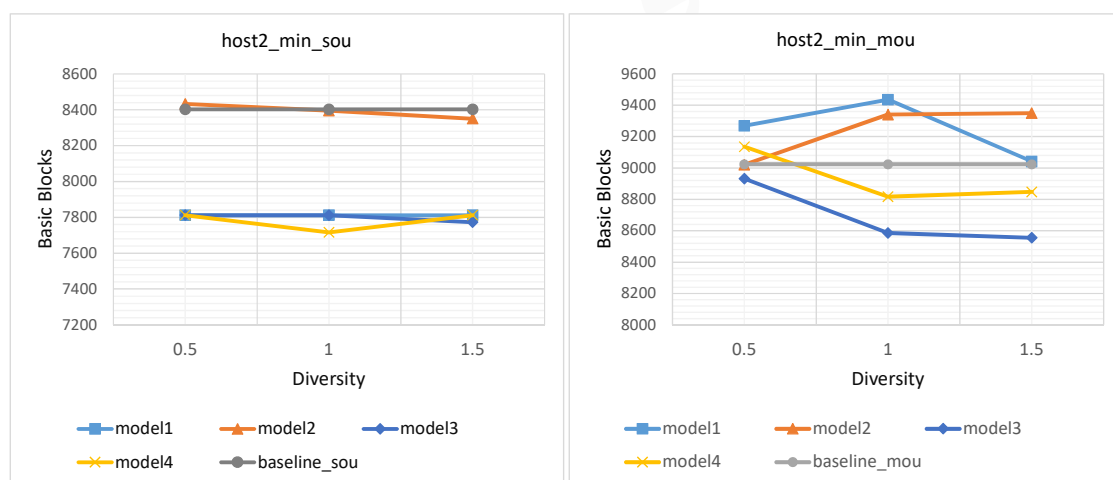
مشاهدات

- پوشش کد داده‌های تولیدی از پوشش کد مبنا در اکثر موارد کمتر است؛ زیرا، اشیای تولید شده به خوش‌شکلی اشیای واقعی نیستند. البته در برخی موارد شاهد افزایش پوشش کد هستیم. از جمله برای host2_min در حالت MOU.
- تولید داده با تنوع ۱ در بیشتر مدل‌ها و روی اکثریت میزبان‌ها پوشش کد بهتری داشته است. یعنی پارامتر تنوع واقعاً سبب بدشکل شدن و بالا رفتن گوناگونی داده‌های تولید شده و در نتیجه پارامتری مؤثر بوده است.
- افزایش تنوع در مدل LSTM دوسویه در حالت کلی سبب افزایش پوشش کد شده است؛ اما در دیگر مدل‌ها خیر.
- تقریباً در همه نمودارها داده‌های تولید شده توسط مدل ۲، پوشش کد بیشینه را نتیجه داده است. یعنی مدل‌های ژرف ساده بهتر از مدل‌های ژرف پیچیده همچون LSTM دوسویه (مدل ۴)، عمل کرده‌اند. به‌همین دلیل مدل ۲ با تنوع تولید داده ۱ پیروز نهایی این سری از آزمایش‌های ما هستند.

^۱ یعنی هر بار یک لیست از ۱۰۰ شیء توسط مدل مولد بازگردانیده می‌شود.

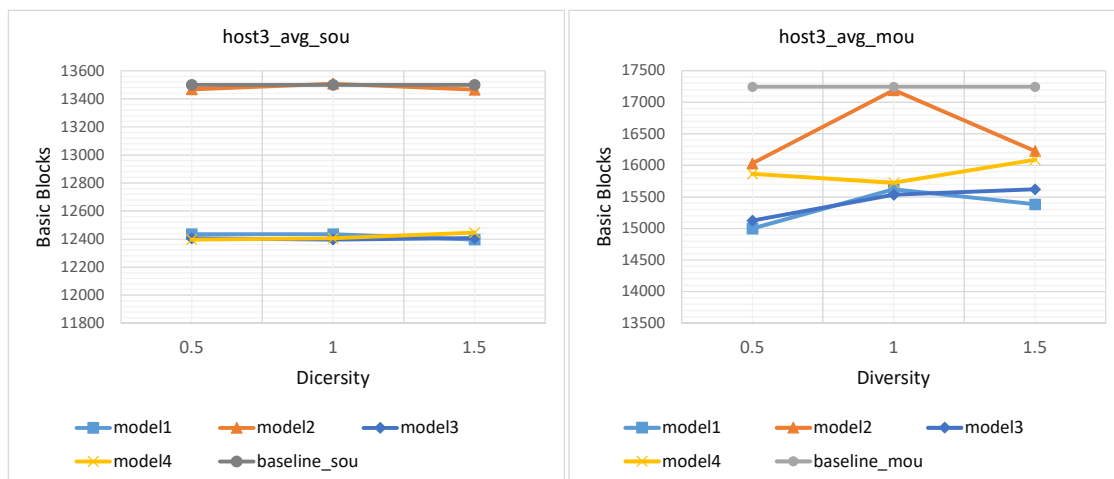


(آ) نمودار تغییرات پوشش کد در تنوع‌های نمونه‌برداری ۰/۵ تا ۱/۵ برای host1_max.

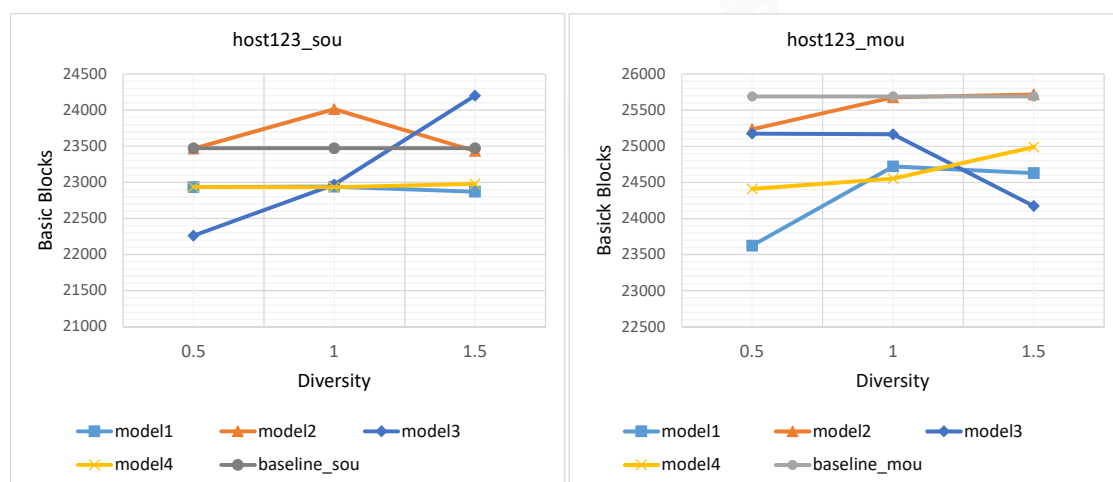


(ب) تغییرات پوشش کد در تنوع‌های نمونه‌برداری ۰/۵ تا ۱/۵ برای host2_min.

شکل ۵-۳: نمودار تغییرات پوشش کد مدل‌های مختلف برحسب تنوع.



(پ) تغییرات پوشش کد در تنوع‌های نمونه‌برداری ۰/۵ تا ۱/۵ برای host3_avg.



(ت) تغییرات پوشش کد در تنوع‌های نمونه‌برداری ۰/۵ تا ۱/۵ برای host123.

شکل ۵-۳: (ادامه). نمودار تغییرات پوشش کد مدل‌های مختلف بر حسب تنوع.

۳-۴-۵ مقایسه با مدل کدگذار-کدگشا

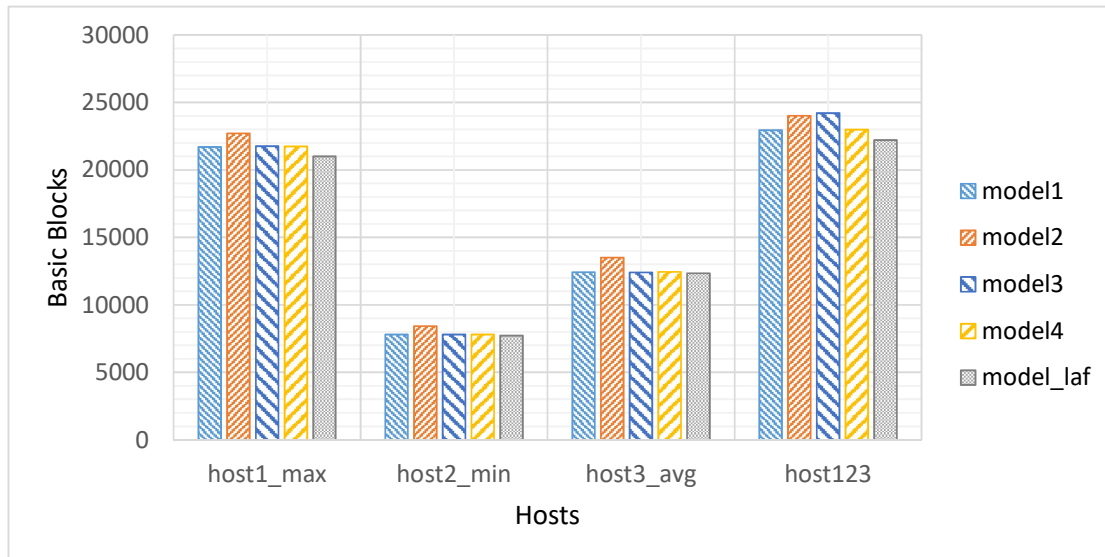
برای مقایسه با [۱۱] ابتدا مدل کدگذار-گشای توضیح داده شده را پیاده‌سازی و آن ۵۰ دوره آموزش دادیم. سپس ۱۰۰۰ فایل PDF را با این مدل تولید و پوشش کد مجموع آنها را اندازه‌گیری کردیم. چون راهبرد نمونه‌برداری برای مدل مذکور نیز بهترین روش گزارش شده است، ما نیز از نمونه‌برداری برای تولید داده‌ها با این مدل استفاده کرده‌ایم. برای مدل‌های چهارگانه خود نیز بالاترین پوشش کد کسب شده در میان همه تنوع‌های آزمایش شده در بخش ۲-۴-۵ را به‌عنوان نماینده پوشش کد برای هر مدل انتخاب کردیم. در نهایت دو حالت SOU و MOU را به‌صورت مجزا مورد ارزیابی و مقایسه قرار دادیم. نتایج در شکل ۴-۵ نشان داده شده است. شکل ۴-۵ آ حالت SOU و شکل ۴-۵ ب حالت MOU را نشان می‌دهد.

مشاهدات

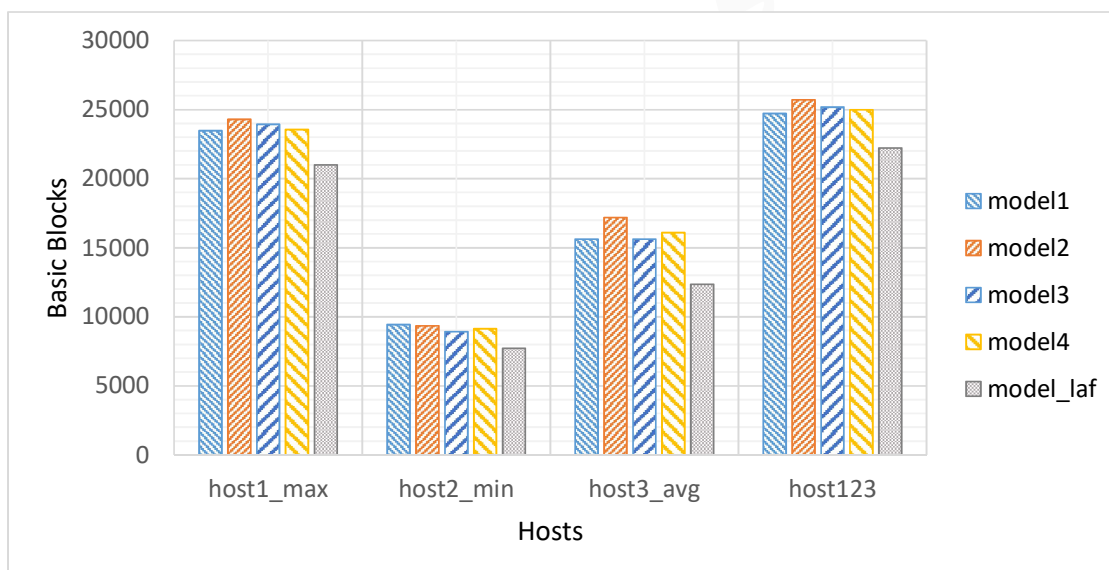
- در حالت SOU و برای $host1_max$ همه مدل‌های پیشنهادی پوشش کد بالاتری داشته‌اند. در همین حالت دو میزبان دیگر اختلاف پوشش کد ناچیز بوده است. اما در حالت اجتماع پوشش کدها یعنی $host123$ مدل‌های پیشنهادی بهتر ظاهر شده‌اند یعنی هر مدل برای هر میزبان مجموعه دستورات متفاوتی را اجرا کرده است.
- در حالت MOU مدل‌های پیشنهادی ما با اختلاف بهتر هستند. این نشان می‌دهد که تغییر بیشتر فایل‌های میزبان به افزایش پوشش کد، در تعداد داده آزمون مساوی، می‌انجامد. بنابراین حالت MOU را برای انجام آزمون فازی توصیه می‌کنیم.
- در هر دو حالت اختلاف بین پوشش کد مدل‌ها برای $host2_min$ کمتر و برای $host1_max$ بیشتر به‌چشم می‌خورد. این امر نشان می‌دهد که انتخاب فایل میزبان در مواردی که نیاز به آن هست، مانند فایل PDF که ساختار پیچیده‌ای دارد و یادگیری کامل آن به‌آسانی محقق نمی‌شود، بسیار حائز اهمیت است و تأثیر چشم‌گیری روی ارتقاء پوشش کد دارد. این انتخاب نبایست تصادفی انجام شود، چون دیدیم افزایش پوشش کد رابطه مستقیمی با پوشش کد فایل میزبان دارد. بنابراین فایل میزبان با پوشش کد بیشتر را برای انجام آزمون فازی توصیه می‌کنیم.

۴-۴-۵ مقایسه در دوره‌های مختلف

یک پارامتر قابل ارزیابی که در جدول ۱-۵ مطرح کردیم و تأثیر آن روی پوشش کد در [۱۱] نیز بررسی شده است تعداد دوره‌های آموزش است. به نظر می‌رسد که پوشش کد ارتباط مستقیمی با تعداد دوره‌های آموزش

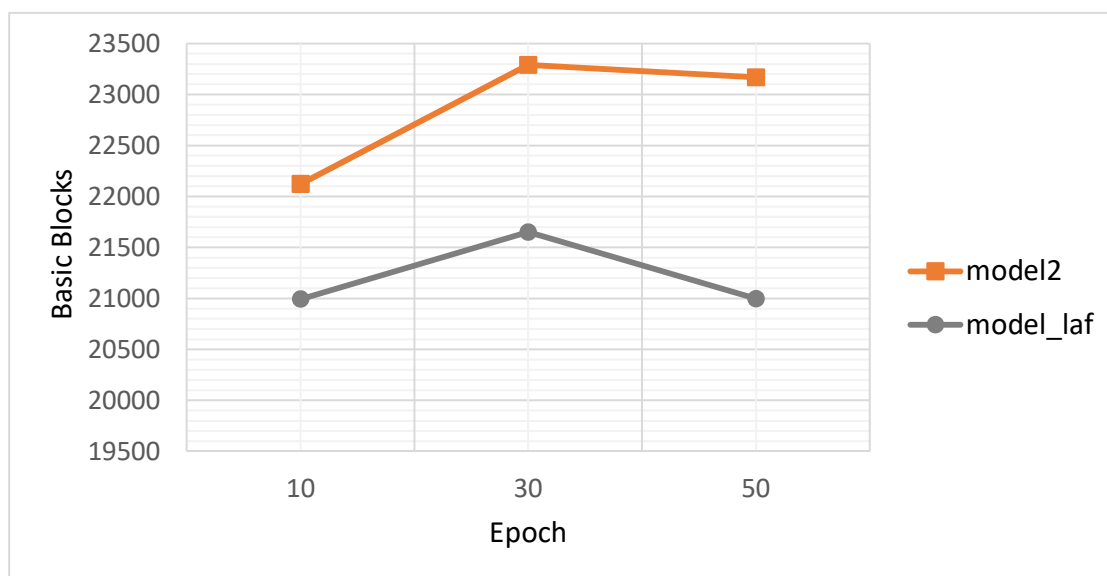


(آ) حالت SOU.



(ب) حالت MOU.

شکل ۴-۵: نمودار پوشش کد مدل‌های مختلف در مقایسه با مدل laf برحسب فایل‌های میزبان.



شکل ۵-۵: نمودار تغییرات پوشش کد برحسب دوره برای مدل‌های ۲ و laf.

مدل ژرف داشته باشد. برای این منظور پوشش کد ۱۰۰۰ فایل تولید شده با استفاده از مدل ۲ را در سه دوره ۱۰، ۳۰ و ۵۰ اندازه‌گیری کردیم. سپس همین کار را برای مدل laf نیز انجام دادیم. برای آن که نتایج قابل اطمینان باشند و اثر پارامترهای تصادفی از بین برود، هر آزمایش را سه مرتبه با سه مجموعه داده مجزا تکرار و میانگین پوشش کدها اندازه گرفتیم. نتایج در شکل ۵-۵ نشان داده شده‌اند.

مشاهدات

- برای هر دو مدل پوشش کد از دوره ۱۰ به ۳۰ افزایش و سپس در دوره ۵۰ کاهش یافته است. یعنی رابطه مستقیم معناداری بین پوشش کد و تعداد دوره‌های آموزش مدل وجود ندارد.
- با بررسی مقدار خطای مدل در هر دوره متوجه شدیم تا زمانی که نرخ کاهش خطای مدل در فرایند آموزش زیاد است، پوشش کد نیز افزایش می‌یابد؛ چراکه ساختار فایل بهتر یادگیری شده و در نتیجه مدل قادر به تولید داده‌های خوش‌شکل‌تری خواهد بود. هنگامی که نرخ کاهش خطا از یک حد آستانه کمتر می‌شود، افزایش پوشش کد نیز ثابت شده یا اندکی کاهش می‌یابد.
- در همه دوره‌های شکل ۵-۵ پوشش کد مدل پیشنهادی از مدل laf بیش‌تر بوده است که نشان می‌دهد این مدل ساختار فایل را بهتر یادگرفته است.

۵-۴-۵ آزمون فازی عصبی

در پنجمین و آخرین آزمایش نرم‌افزار MuPDF را با استفاده از فازر پیشنهادی در بخش ۵-۴، مورد آزمون فازی قرار دادیم. برای این منظور الگوریتم‌های فاز عصبی داده و فاز عصبی فراداده (الگوریتم‌های ۲-۴ و ۳-۴) را که در بخش ۳-۴ معرفی کرده بودیم، پیاده‌سازی و با استفاده از هریک از آنها تعداد ۱۰۰۰۰ فایل PDF را ایجاد کردیم. تنظیمات پارامترهای ورودی و ثوابت موجود در الگوریتم‌های پیشنهادی به قرار جدول ۵-۵ است. ستون آخر جدول ۵-۵، همچنین مقادیر مجاز برای هریک از ورودی‌ها و ثوابت داده شده را نشان می‌دهد. حداکثر طول یک شیء داده‌ای PDF را در بازه متغیر ۴۵۰ تا ۵۵۰ انتخاب کرده‌ایم؛ زیرا، به طور میانگین طول اشیای استخراج شده مجموعه‌های آموزش و آزمون، در این بازه قرار دارد. در آزمایش‌های آتی خود در نظر داریم تا آزمون فازی را با مقادیر متنوع انتخابی از مجموعه‌های داده شده انجام دهیم و بدین ترتیب قادر خواهیم بود تا اثر هریک از این پارامترها را به طور دقیق‌تری بررسی کنیم.

ما همچنین نسخه‌هایی از الگوریتم‌های ۱-۳ [۲] و ۴-۳ [۱۱] را پیاده‌سازی و آزمون فازی را با تولید داده آزمون از طریق این الگوریتم‌ها نیز انجام دادیم. برای حالت تصادفی (الگوریتم ۱-۳) از ابزار FileFuzz که یک فازر تصادفی تحت سیستم عامل ویندوز است و در بخش ۱-۲-۲ معرفی شد، استفاده کردیم. متأسفانه در روش یادگیری و فاز [۱۱]، تقریباً اکثر ابرپارامترهای مورد نیاز در هنگام آزمایش‌ها نامشخص هستند و نویسندگان اشاره‌ای به مقادیر آنها نکرده‌اند. برای این هریک از این ابرپارامترها نظیر d ، ما همان مقدار استفاده شده در الگوریتم‌های روش پیشنهادی را استفاده کردیم، تا بدین ترتیب شرایط یکسانی بر آزمایش‌ها حاصل باشد. در همه آزمایش‌ها، `host1_max` به عنوان میزبان^۱ استفاده شد. نتایج حاصل از پوشش کد روش‌های مختلف در جدول ۵-۶ آمده است. همچنین جدول ۷-۵ میزان بهبود در پوشش کد روش پیشنهادی در مقایسه با کارهای مرتبط را نشان می‌دهد.

مشاهدات

- الگوریتم فاز عصبی فراداده به پوشش کد کمتری دست پیدا کرده است، زیرا همان‌طور که انتظار می‌رود دستکاری بخش کوچکی در از قالب فایل ممکن است آن را کاملاً نامعتبر کند و در مرحله اولیه تجزیه توسط تجزیه‌گر رد شود. بنابراین الگوریتم در رسیدن به هدف خود یعنی فاز فراداده موفق بوده است.

- هر دو الگوریتم فاز عصبی داده و فاز عصبی فراداده پوشش کد بهتری نسبت به SampleFuzz داشته‌اند که نشان از عملکرد بهتر مدل‌های پیشنهادی در یادگیری ساختار فایل و عملکرد بهتر الگوریتم‌های

^۱ در تولید داده آزمون به روش تصادفی و روش مبتنی بر جابه‌جایی، به جای میزبان، به یک یا تعدادی دانه اولیه نیاز داریم که در اینجا از همان `host1_max` استفاده گردید.

جدول ۵-۵: مقادیر مجاز و مقادیر داده شده به ثوابت و ورودی‌های الگوریتم‌های فازی عصبی داده و فاز عصبی فرا داده در هنگام آزمون فازی قالب فایل PDF

پارامتر ورودی / ثابت	مقادیر استفاده شده	مقادیر مجاز
Learnt model M	۲	$1, 2, 3, 4, laf$
Sequence prefix P	انتخاب از مجموعه آزمون	ثابت رشته‌ای
Diversity D	۱	$(0, +\infty)$
Fuzzing rate FR	۰/۱۰	$(0, 1]$
End token ET	endobj	ثابت رشته‌ای
Binary token BT	stream	ثابت رشته‌ای
(a, b)	$(450, 550)$	$(len(P), +\infty)$
α (DataNeuralFuzz)	۰/۵۰	$(0, 1)$
β (MetadataNeuralFuzz)	۰/۹۰	$(0, 1)$

جدول ۵-۶: نتایج پوشش کد حاصل از آزمون فازی نرم‌افزار MuPDF با الگوریتم‌های تولید داده آزمون مختلف.

الگوریتم تولید داده آزمون / معیار	پوشش بلوک پایه درصد	پوشش خط‌کد درصد
DataNeuralFuzz	۲۳۷۱۹	۱۹/۳۶
MetadataNeuralFuzz	۲۲۵۸۳	۱۸/۴۳
SampleFuzz [۱۱]	۲۰۹۵۷	۱۷/۱۰
RandomFuzz (FileFuzz) [۲]	۷۵۶۳	۶/۱۷

جدول ۵-۷: میزان بهبود پوشش کد ابزار MuPDF توسط الگوریتم‌های روش پیشنهادی در مقایسه با کارهای مرتبط. اعداد داخل جدول به صورت درصد هستند. مقادیر هر خانه برابر اختلاف پوشش کد حاصله از الگوریتم‌های ستون و سطر مربوط به آن مقدار است.

روش پیشنهادی		
روش‌های موجود	DataNeuralFuzz	MetaDataNeuralFuzz
SampleFuzz [۱۱]	+۲/۲۶	+۱/۸۳
AFL [۲۰]	+۷/۷۳	+۶/۸۰
AugmentedAFL [۲۰]	+۷/۵۶	+۶/۶۳
RandomFuzz (FileFuzz) [۲]	+۱۳/۱۹	+۱۲/۲۶

پیشنهادی در فاز کردن این فایل‌ها دارد. پوشش کد حاصل شده در این آزمایش‌ها همچنین از پوشش کد گزارش شده در [۲۰] که حاصل از آزمون فازی MuPDF با AFL و AFL افزوده بیشتر است. اعداد مربوط به آنها قبلاً در بخش ۱-۲-۳ ذکر شدند. جدول ۷-۵ در این بخش نیز اختلاف این مقادیر را با مقادیر حاصل از پوشش کد روش پیشنهادی نشان داده است.

- به‌وضوح می‌توان برتری روش‌های هوشمند تولید داده آزمون را نسبت به روش‌های تصادفی مشاهده کرد. همانطور که در ابتدای این پایان‌نامه گفتیم روش‌های تصادفی در ساختارهای پیچیده پوشش کد بسیار کمی را نتیجه می‌دهند. پوشش کد الگوریتم فاز عصبی داده بیش از ۳ برابر الگوریتم فاز تصادفی است.

- باوجود استفاده از هوش مصنوعی و الگوریتم‌های هوشمند در تولید داده آزمون همچنان شاهد پوشش کد پایینی (کمتر از ۲۰ درصد) در پایان عملیات آزمون هستیم. به‌نظر می‌رسد رسیدن به پوشش کد بالا در ساختارهای پیچیده، با آزمون فازی جعبه سیاه کار دشواری باشد. از طرفی فنون آزمون جعبه سفید مانند روش‌های اجرای نمادین نیز روی این ساختارها به‌علت پیچیدگی بالا و وجود محدودیت‌های فراوان سخت، زمان‌بر و تا حد زیادی غیرممکن است و کماکان آزمون فازی مؤثرتر بوده است.

خطاها و آسیب‌پذیری‌های شناسایی شده

در بررسی گزارش‌های تولید شده توسط Application Verifier پس از هر آزمون، هیچ‌گونه خطایی مشاهده نکردیم. باتوجه به‌اینکه ما نسخه نهایی نرم‌افزار MuPDF را مورد آزمایش قرار دادیم، تصور می‌شود که بیشتر خطاهای آن در نسخه‌های آزمایشی برطرف شده باشد و در نتیجه پیدا کردن خطای جدید سخت خواهد بود. از طرفی MuPDF نرم‌افزاری تحت توسعه فعال و جامعه توسعه‌دهنده و کاربری بزرگی است که سبب می‌شود تا از کیفیت خوبی برخوردار باشد. با این حال الگوریتم فاز عصبی داده چندین مورد استفاده از توابع ناامن را شناسایی کرد که Application Verifier آنها را در قالب هشدارهای امنیتی اعلام کرده است.

پیش از آزمایش‌های اصلی ما فازر و Application Verifier روی نرم‌افزارهای کوچکی که خطای آنها مشخص بود اجرا کردیم. به‌نظر می‌رسد Application Verifier روی ویندوز ۱۰ نسخه ۶۴ بیتی قادر به تشخیص خطاهای حافظه برنامه‌های ۳۲ بیتی نیست؛ زیرا موفق به شناسایی این خطاها نشدیم. برای برنامه‌های ۶۴ بیتی اما این مشکل وجود نداشت. لذا ما هر دو نسخه ۳۲ و ۶۴ بیتی MuPDF را مورد آزمون قرار دادیم. فازر، برنامه PDF خوان را با داده آزمون ورودی باز و بعد از ۱۰ ثانیه آن را بسته و برای تزریق داده آزمون بعدی اقدام می‌کند. همزمان پوشش کد و گزارش وضعیت حافظه نیز ثبت و ذخیره می‌شوند. بنابراین زمان

آزمون برای هر مجموعه ۱۰۰۰۰ فایلی حدود ۲۸ ساعت به طول انجامید^۱. برای همه نرم افزارهای دیگر آزمون به این شیوه قابل انجام خواهد بود. گفتیم که آزمون فازی نوعی آزمون فشار است. در نظر داریم تا آزمون فازی با این روش را با استفاده از ۱۰۰ هزار فایل تولید شده و حتی بیشتر انجام دهیم. در این صورت امکان سقوط برنامه MuPDF افزایش می یابد.

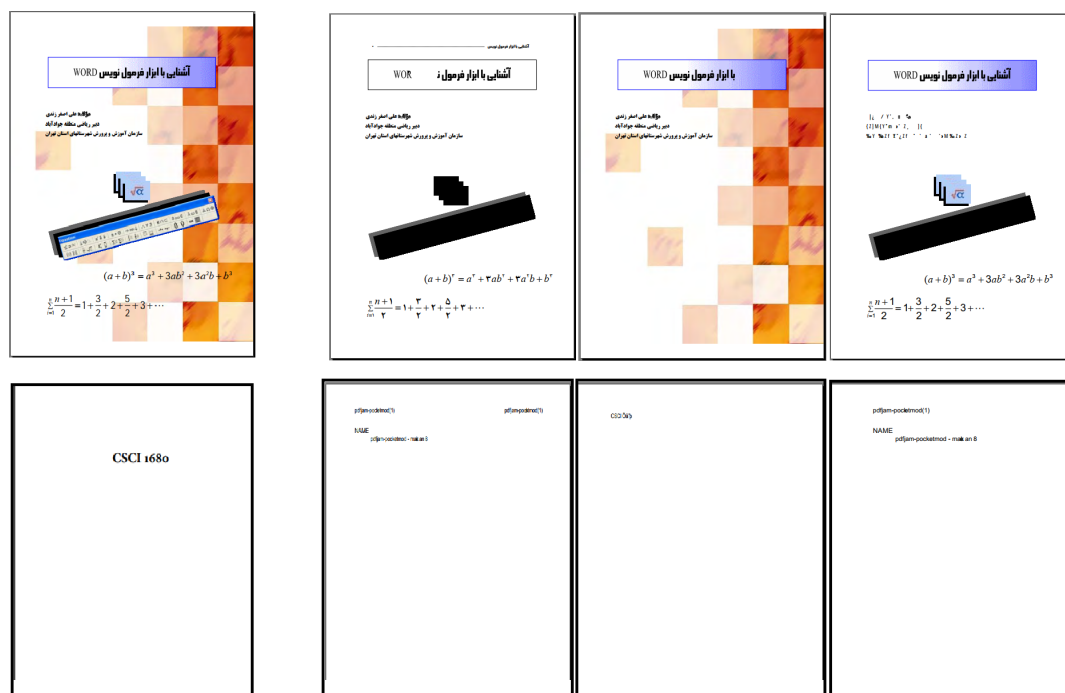
تنوع داده های تولید شده

مدل های ما در روش پیشنهادی قادر به تولید داده های متنوع و به صورت کنترل شده بدشکل هستند. الگوریتم فاز عصبی داده به خوبی در تغییر محتویات فایل های PDF عمل کرده و می تواند فایلی با داده های جدید بدون تغییر در ساختار ایجاد کند. شکل ۵-۶ نمونه ای از داده های تولید شده توسط الگوریتم فاز عصبی داده را نشان می دهد. همان طور که در این شکل پیداست فایل های PDF در عین معتبر بودن حاوی داده های جدیدی هستند که با مقادیر مرزی در فرایند تولید داده آزمون جایگذاری شده اند. به دلیل معتبر بودن تعداد بیشتری از فایل های تولید شده توسط الگوریتم فاز عصبی داده پوشش کد این الگوریتم از پوشش کد الگوریتم فاز عصبی فراداده بیشتر است. اما هر دو الگوریتم مورد نیاز هستند؛ زیرا، هر کدام مرحله مجزایی از مراحل پردازش فایل در کد برنامه را هدف آزمون قرار می دهند.

۵-۵ خلاصه

در این فصل، ابتدا پارامترهای قابل بررسی و مؤثر در روش های تولید داده آزمون مبتنی بر یادگیری ژرف را معرفی و سپس روش پیشنهادی خود را که در فصل ۴ مطرح کرده بودیم، روی نرم افزار MuPDF و برای مهمترین این پارامترها مورد آزمایش و ارزیابی قرار دادیم. نتایج در حالت کلی حاکی از بهبود میزان پوشش کد در آزمون فازی MuPDF هنگام استفاده از مدل های و الگوریتم های تولید داده آزمون پیشنهاد شده است. به خصوص نسبت به تعدادی از روش های قبلی از جمله [۱۱] که مشابه ترین کار مرتبط بود، آمار و ارقام بهتری هم در دقت مدل ها و هم در پوشش کد مشاهده می شود. علاوه بر این نتیجه گیری کلی آزمایش های مختلف ما چندین واقعیت تجربی دیگر را مشخص کرد که مهمترین آنها عبارتند از:

^۱ Application Verifier و ابزارگذاری کد سربارهایی را به زمان هر بار اجرای برنامه اضافه می کنند. همچنین فاصله زمانی کوتاهی بین تزریق داده های آزمون متوالی در نظر گرفته شد؛ زیرا، در مواردی سقوط برنامه ناشی از خطای دیگر بخش های سیستم است. با لحاظ این زمان ها از سقوط های این چنینی با اطمینان زیادی جلوگیری می نمایم.



شکل ۵-۶: نمونه‌ای از داده‌های آزمون متنوع تولید شده توسط الگوریتم فاز عصبی داده در فرایند آزمون فازی قالب فایل PDF. در هر ردیف، سمت چپ‌ترین فایل، فایل میزبان را نشان می‌دهد و ۳ فایل سمت راست فایل‌های حاصل از مدل مولد هستند. تغییر در محتوا به وضوح قابل مشاهده است.

- LSTM دوسویه به عنوان مدل زبانی می‌تواند به دقت بیشتر و خطای کمتری روی مجموعه داده دست پیدا کند. با این حال مدل‌های ژرف ساده مانند LSTM یک‌سویه بدون لایه‌های Dropout توانستند مدل‌های پیچیده‌تری مثل LSTM دوسویه را در آزمون فازی شکست دهند. نتیجه‌گیری مشابهی در [۲۰] بیان شده است.
- فایل‌های PDF ای که پوشش کد بیشتری دارند، هنگام تغییر نحوه بروزرسانی افزایشی نیز، اختلاف پوشش کد بیشتری نسبت به دیگر فایل‌ها فراهم می‌کنند.
- تنوع تولید داده ۱ برای بیشتر مدل‌ها منجر به پوشش کد بهتری می‌شود.
- افزایش دوره‌های آموزش لزوماً به افزایش پوشش کد نمی‌انجامد اما تا زمانی که خطای مدل‌ها را به نحو خوبی کاهش دهد، پوشش کد را افزایش می‌دهد.
- روش‌های ترکیبی تولید داده آزمون مانند الگوریتم‌های فاز عصبی داده و فاز عصبی فراداده، که بخش‌های دودویی را نیز در فرایند آزمون فازی شرکت می‌دهند منجر به پوشش کد بیشتری می‌شوند.

• آزمون فازی با روش‌های تولید داده هوشمند، آزمون فازی تصادفی را همواره شکست می‌دهد ولی هنوز هم روی ساختارهای خیلی پیچیده به پوشش کد آن‌چنان بالایی منتهی نمی‌گردد، این امر ارزشمند بودن کار بر روی روش‌های تولید داده آزمون در آینده را نشان می‌دهد. البته دلیل پوشش کد پایین برای نرم‌افزار MuPDF بیشتر آن است که این نرم‌افزار قالب‌های فایل دیگری غیر از PDF را پشتیبانی می‌کند. لذا اجرای آن با فایل‌های PDF تنها کدهای مربوط به تجزیه و پرداخت PDF را سبب می‌شود. بنابراین بایستی قالب‌های دیگر را نیز در فرایند آزمون شرکت داد که بدون شک منجر به افزایش پوشش کد می‌شود.

اثر تعدادی از پارامترهای دیگر مطرح شده در جدول ۵-۱ مانند راهبردهای تولید داده از مدل، در این آزمایش‌ها بررسی نشدند که جای دارد آزمایش‌هایی برای آنها نیز طراحی کرد. با این حال باتوجه به محدودیت‌های که برای هریک از دیگر راهبردها مثل راهبرد حریصانه برشمرده شد، راهبرد نمونه‌برداری مناسب‌ترین راهبرد به نظر می‌رسد.

فصل ۶

نتیجه‌گیری و کارهای آتی

«ما ممکن است امیدوار باشیم که ماشین‌ها در نهایت در همه زمینه‌های هوشمند با انسان رقابت کنند، اما بهترین زمینه برای شروع کدام است؟!»

✠ آلن تورینگ

۱-۶ نتیجه‌گیری

دکتر اندرو ان‌جی^۱ هوش مصنوعی را یک الکتریسیته جدید می‌نامد که می‌تواند تحول بزرگ بعدی را در صنعت رقم بزند. یادگیری ژرف و نگاه متفاوت آن به حل مسئله تا همین الان این تحول بزرگ را در وظیفه‌هایی مانند پردازش تصویر، پرازش صوت، پردازش متن و ترجمه ماشینی رقم زده است. سال ۲۰۱۷ آغاز استفاده از یادگیری ماشینی در آزمون فازی و تولید داده آزمون بود [۲۰، ۱۰]. پژوهشگران مایکروسافت برای نخستین بار از این فنون در آزمون فازی استفاده کردند. در این پایان‌نامه ما از مدل‌های یادگیری ژرف برای یادگیری ساختار فایل‌های پیچیده و سپس تولید داده آزمون جدید به منظور استفاده در فرایند آزمون فازی استفاده کردیم. به طور خاص هر فایل را می‌توان نمونه‌ای مشتق شده از زبان یا گرامر ساختار آن دانست. براین اساس ما با استفاده از شبکه‌های عصبی مکرر اقدام به ایجاد یک مدل زبانی عصبی برای هر ساختار فایلی می‌کنیم که یک توزیع احتمالی از چگونگی وقوع نشانه‌ها در یک فایل را با آموزش روی یک مجموعه داده تخمین می‌زند. سپس این مدل زبانی را برای تولید فایل‌های جدید به کار می‌بندیم. یک ویژگی مفید امکان تشخیص و تمایز بین داده و

^۱ Andrew Ng (<http://www.andrewng.org/>)

فراداده در یک فایل با استفاده از چنین مدل‌هایی است در نتیجه می‌توان جابه‌جایی آزمون فازی را با استناد به تفکیک داده و فراداده به‌نحو هوشمندتری انجام داد. دو الگوریتم پیشنهادی در این راستا نشان دادند چنین روشی قادر به اجرای بخش‌های بیشتری از کد یک SUT است و پوشش کد بیشتری را رقم می‌زند که در نتیجه امکان شناسایی خطا نیز افزایش خواهد یافت.

در فصل ۱، مسائلی را در باب سخت بودن تولید داده آزمون برای رسیدن به پوشش کد بالا در آزمون فازی قالب‌های فایل با ساختار پیچیده مثل PDF مطرح کردیم و دیدیم که چنانچه بتوان با استفاده از گرامر ورودی داده‌های آزمون را تولید کرد، تعداد داده‌هایی که در مراحل اولیه توسط کدهای مدیریت استثنای تجزیه‌گر رد می‌شوند کاهش یافته و قادر به نفوذ به مسیرهای عمیق‌تر برنامه خواهیم بود. با این ایده و براساس مطالعات اولیه و کارهای قبلی که در فصل‌های ۲ و ۳ به آنها اشاره کردیم، در فصل ۴ یک روش یادگیری ساختار فایل را پیشنهاد دادیم که تولید داده آزمون مبتنی بر گرامر را خودکار می‌کند. در فصل ۵ هم ابتدا پارامترهایی را که در تولید خودکار داده آزمون از روی مدل‌های یادگیری ژرف، نقش دارند شناسایی و سپس آزمایش‌هایی را برای بررسی و مقایسه تأثیر هریک از این پارامترهای طرح‌ریزی، پیاده‌سازی و اجرا نمودیم.

یک نتیجه‌گیری مهم که آزمایش‌های ما نشان می‌دهند این است که افزایش ظرفیت محاسباتی و پیچیدگی شبکه‌های عصبی ژرف که در نتیجه قدرت حل مسئله یک شبکه را بالا می‌برد لزوماً منجر به نتایج بهتری در آزمون فازی نخواهد شد و مشاهده کردیم که مدل‌های ساده مدل‌های پیچیده‌تر را در پوشش کد SUT شکست می‌دهند. این مدل‌ها طبیعتاً به‌زمان کمتری برای آموزش نیاز داشته و تولید داده با استفاده از آنها نیز سریع‌تر خواهد بود. بنابراین از همه نظر مقرون به صرفه هستند.

روش پیشنهادی در این پایان‌نامه، در عین مزایایی که برای آن برشمرده شد محدودیت‌ها و نقاط ضعفی دارد. بزرگترین محدودیت را می‌توان عدم اطلاع مدل مولد از وضعیت SUT دانست. به عبارت بهتر مدل مولد در غیاب SUT و تنها روی مجموعه داده‌های یک قالب فایل آموزش دیده و سپس اقدام به تولید داده‌های آزمون می‌کند. فازر پیشنهادی نیز فاقد یک حلقه بازخورد برای دریافت اطلاعات زمان اجرای SUT است. این در حالی است که همواره بازخوردهای اجرای SUT می‌تواند حاوی اطلاعات خوبی برای پیش‌برد ادامه فرایند آزمون در اختیار فازر قرار دهد. این محدودیت می‌تواند به‌عنوان کارآتی مورد بررسی قرار گیرد. مشکل بعدی که البته مختص به روش پیشنهادی ما نبوده و هر روش یادگیری ماشینی در این زمینه با آن مواجه است نیاز به تعدادی زیادی فایل به‌عنوان مجموعه داده است. برای قالب‌های مشهور فایل مانند PDF که مورد مطالعاتی این پایان‌نامه بود، این مشکل پُررنگ نیست اما اگر برای یک قالب خاص امکان تهیه مجموعه داده بزرگی نباشد، استفاده از این روش تقریباً غیر ممکن خواهد بود.

در هر حال تلاش در راستای موضوع این پایان‌نامه صرف نظر از نتایج تجربی آن به دلیل پیوند دو شاخه به‌ظاهر کمتر مرتبط در علم کامپیوتر یعنی یادگیری ژرف و آزمون فازی ارزشمند به نظر می‌رسد. به‌ویژه که

طراحی و آموزش شبکه‌های عصبی ژرف کاری مهیج بوده و این مسیر پژوهشی نیز در ابتدای راه خود است. در این بین دو حوزه مذکور که این پایان‌نامه بر آنها بیان شده، نیز هرکدام در حالت کلی دارای مزایا و معایبی هستند که ممکن است برخی از آنها را قبلاً هم ذکر کرده باشیم، با این حال در پایان این بخش نگاهی کوتاه به مزایا، معایب و آینده هریک از این دو حوزه خواهیم داشت؛ زیرا می‌توانند در تعیین روند پژوهش‌های آتی مؤثر واقع شوند.

۶-۱-۱ مزایا و معایب یادگیری ژرف

یادگیری ژرف و به تبع آن شبکه‌های عصبی ژرف در انجام وظایف ساده برای انسان، سخت برای ماشین بسیار موفق ظاهر شده‌اند. با توجه به افزایش قدرت محاسبات انجام حجم وسیعی از محاسبات در مسائل پیچیده، ارزان‌تر از نوشتن یک الگوریتم خاص می‌باشد به نحوی که در آینده شاهد افزایش ظرفیت‌های سخت‌افزاری برای توسعه چنین مدل‌هایی در مقیاس‌های بسیار بزرگ خواهیم بود. دو ویژگی بسیار مهم این شبکه‌ها عبارتند از تعمیم‌پذیری^۱ و تطبیق‌پذیری^۲ [۳۸]. تعمیم‌پذیری بدین معنی است که در صورت آموزش صحیح، شبکه برای ورودی‌های جدید نیز درست کار خواهد کرد. تطبیق‌پذیری یعنی در صورتی که داده‌ها تغییر کند، شبکه هم توانایی تغییر خواهد داشت. یعنی یک مدل می‌تواند برای حل خانواده‌ای از مسائل مشابه طراحی و ساخته شود و هر بار با داده‌های مختلفی آموزش ببیند.

شبکه‌های عصبی سراسر فایده و نوش‌داروی حوزه محاسبات جدید نیستند و در عین حال معایبی دارند. از جمله اینکه آموزش آنها سخت و بسیار مستعد خطا است. در واقع دقت نتایج بستگی زیادی به مجموعه آموزش دارد؛ به نحوی که یک مجموعه آموزش ضعیف (کوچک یا نادرست) عملاً شبکه غیرقابل استفاده‌ای را نتیجه می‌دهد. دیگر آنکه قوانین مشخصی برای طراحی یک شبکه جهت کاربردی خاص وجود ندارد. به عبارت بهتر تعیین ابرپارامترها مسئله تصمیم‌ناپذیر است؛ یعنی یک نمی‌توان به صورت الگوریتمی بهترین مجموعه ابرپارامتر را برای یک شبکه در یک وظیفه خاص تعیین کرد. این کار معمولاً با سعی و خطا انجام می‌شود و بالاخره اینکه نمی‌توان به فیزیک یا قانون حاکم بر مسئله حل‌شده توسط شبکه پی برد و تنها با مشتی اعداد سروکار خواهیم داشت که روی یک مجموعه یک هدف خواسته شده را بهینه کرده‌اند.

^۱Generalization

^۲Adaptive

۶-۱-۲ مزایا و معایب آزمون فازی

آزمون فازی چندین سودمندی دارد. در درجه نخست سادگی و راحتی خودکارسازی فرایند شرح داده شده است. به همین دلیل، فازرهای بسیاری توسعه داده شده است. برای استفاده از فازرهای موجود تنها انتخاب SUT و فراهم ساختن تعدادی داده آزمون اولیه یا قالب ورودی لازم است. فازر می‌تواند به صورت یک پردازش پس‌زمینه^۱ و بدون دخالت اضافی کاربر، در یک حلقه بی‌نهایت، به مدت طولانی اجرا شود. برای ساخت یک فازر جدید نیز کافی است پیمانه‌های شکل ۲-۲، توسعه داده شده و در کنار هم قرار گیرند.

آزمون فازی همچنین کاستی‌های آزمون معمولی که به صورت دستی صورت می‌پذیرد را جبران می‌کند. آزمون‌های نوشته شده به صورت دستی تا حدودی تمایل به پیش‌دانسته‌های ذهنی فرد آزمون‌گر در مورد کد دارند. آزمون فازی از انحراف یاد شده مستثنی است و می‌تواند ورودی‌های بدشکل را به نحوی تولید کند که پیش از آن هیچگاه، به ذهن فرد یا افراد آزمون‌گر نرسیده است. با گذشت سه دهه از ابداع آزمون فازی این روش جایگاه ویژه‌ای در صنعت و نیز در پژوهش یافته است و به همین دلیل کار بر روی آن ارزشمند و اثر بخش است.

در حالی که آزمون فازی در پیدا کردن خطاهای فساد حافظه، بسیار خوب عمل می‌کند، خطاهای پیچیده‌تر مانند خطاهای منطقی به ندرت توسط این آزمون قابل شناسایی هستند؛ چرا که آشکارسازی آنها نیازمند به اتمام رسیدن اجرای برنامه بدون خطای حافظه و داشتن سروش^۲ آزمون است. افزون بر این تحریک برخی خطاهای حافظه برای فازرها بسیار سخت خواهد بود. مسئله انفجار مسیر که ناشی از وجود حلقه‌های تکرار و شرط‌های تودرتو است مانع از اجرای نمادین برنامه‌های پیچیده می‌شود. در نتیجه به‌کارگیری اجرای نمادین در آزمون فازی جعبه سفید در بسیاری موارد ممکن نیست و این مسئله به پوشش کد پایین و خوب آزمون نشدن برنامه می‌انجامد.

آزمون فازی ممکن است برای اهداف سوء مورد استفاده قرار گیرد. یعنی شناسایی آسیب‌پذیری‌ها توسط فرد مهاجم و بهره‌برداری از آنها جهت حمله به یک سیستم نرم‌افزاری. در نتیجه روش‌هایی برای مقابله با امکان آزمون‌پذیری یک نرم‌افزار مطرح شده‌اند که می‌توان آنها را در طبقه روش‌های ضد مهندسی معکوس دانست. از جمله این روش‌ها می‌توان به حوزه‌ای جدید که اخیراً تحت عنوان ضد فازینگ^۳ مطرح شده است، اشاره کرد که در آن اقداماتی برای کندسازی یا به‌کلی مانع شدن کشف خطاها توسط آزمون فازی انجام می‌گیرد. کاهش کارآمدی و پوشش خطاها دو رویکرد از میان رویکردهای موجود در این زمینه هستند. دیگر فنون مقابله با مهندسی معکوس نظیر مبهم‌سازی کد را نیز می‌توان در ضد فازینگ به‌کار گرفت. چنانچه این رویکردها

^۱ Background Process^۲ Oracle^۳ Anti-fuzzing

به بلوغ خوبی برسند و توسعه‌دهندگان آنها را در برنامه‌های خود لحاظ نمایند، در آینده بایستی به دنبال فنون جایگزین برای روش آزمون فازی جعبه خاکستری و به‌طور کلی آزمون فازی باشیم. البته استفاده از این فن آزمون در بین توسعه‌دهندگان و پیش از انتشار نسخه نهایی نرم‌افزار همچنان در صنعت ادامه خواهد یافت.

۶-۲ نوآوری‌ها

مجموعه نوآوری‌ها، دستاوردها و محصولات کار پژوهشی ما در قالب این پایان‌نامه عبارت است از:

۱. خودکارسازی فرایند یادگیری ساختار فایل ورودی و تولید داده‌های آزمون برای آزمون فازی قالب فایل با به‌کارگیری روش‌های جدید یادگیری ژرف.

۲. بهبود پوشش کد SUT با به‌کار بست روشی ترکیبی یعنی تولید فراداده و داده‌های متنی، مبتنی بر گرامر و سپس تزریق داده‌های دودویی مبتنی بر جابه‌جایی تصادفی در مکان‌های تعیین شده.

۳. شناسایی و استخراج پارامترهای مؤثر در تولید داده آزمون به‌روش یادگیری ژرف و ارزیابی تأثیر آنها با ارایه مجموعه‌ای از آزمایش‌های کنترل شده.

۴. معرفی یک مجموعه دانه اولیه و یک مجموعه داده آزمون برای آزمون فازی قالب فایل PDF تحت یک پیکره از فایل‌های PDF به‌همراه پوشش کدهای آنها.

۵. طراحی و پیاده‌سازی یک فازر قالب فایل ساده با معماری کاملاً پیمانه‌ای و قابل حمل، مجهز به پیمانه تولید خودکار داده‌های آزمون با استفاده از مدل‌های مولد.

به‌طور خلاصه در این پایان‌نامه یک مجموعه داده، چهار مدل مولد، دو الگوریتم فاز و یک فازر قالب فایل پیشنهاد و معرفی گردید که همه آنها در قالب یک بسته نرم‌افزاری تحت عنوان IUST Deep Fuzz منتشر شده‌اند. از این محصول می‌توان در عمل برای آزمون فازی و شناسایی خطاهای نرم‌افزارهایی با ورودی فایل استفاده کرد. تنظیمات کنونی بر روی قالب فایل PDF تعیین شده است اما به‌راحتی قابل تغییر است. نرم‌افزارهای با ورودی فایل علاوه بر PDF خوان‌ها، شامل خانواده وسیع کامپایلرها، مرورگرهای وب، محیط‌های توسعه مجتمع و غیره می‌شوند، که همه آنها در طبقه برنامه‌های کاربردی و بسیار مهم قرار می‌گیرند. تقریباً در همه موارد مذکور، داده کافی برای ایجاد مدل مولد وجود دارد. به‌عنوان مثال در آزمون مرورگرها، ایجاد یک مدل مولد برای تولید فایل‌های ترکیبی HTML، CSS و JavaScript داده آزمون بهتری نسبت به تولید تنها یکی از این فایل‌ها [۴۸] فراهم می‌کند و از طرفی برای هر سه قالب فایل نام‌برده مجموعه داده به فراوانی یافت می‌شود.

۳-۶ ملاحظات اعتبارسنجی

در این پایان‌نامه بر بهبود معیار پوشش کد در آزمون فازی، به‌طور مکرر تأکید ورزیدیم. منظور از پوشش کد در اینجا پوشش دستور (یا در حالت کلی‌تر پوشش بلوک پایه است). در مواردی به پوشش مسیر نیز اشاره کردیم و به‌عنوان مثال در مورد پوشش مسیرهای اجرایی عمیق برنامه صحبت به میان آوردیم. ذکر این نکته ضروری است که معیارهای پوشش دستور و پوشش مسیر متفاوت هستند و همچنان‌که در فصل ۲ دیدیم، پوشش مسیر معیار سخت‌گیرانه‌تری نسبت به پوشش دستور است؛ یعنی ممکن است تمامی دستورات برنامه در یک یا چند اجرا، حداقل یک‌بار اجرا شوند ولی لزوماً تمامی مسیرهای اجرایی پوشش داده نشوند.

یک برنامه با یک دستور *if* ساده را در نظر بگیرید. اگر برنامه در حالتی اجرا شود که حاصل ارزیابی عبارت شرطی درست شود، همه دستورات (همچنین همه بلوک‌های پایه) برنامه اجرا می‌شوند. در این حالت پوشش بلوک پایه معادل ۱۰۰ درصد خواهد بود. در حالی که این برنامه در بدوی‌ترین شکل خود، دو مسیر اجرایی دارد: یک مسیر که از بدنه دستور شرطی *if* عبور می‌کند و مسیر دیگر که وارد بدنه دستور *if* نمی‌شود. برای سناریوی اجرای ذکر شده، پوشش مسیر ۵۰ درصد است. بنابراین تفاوت معناداری میان پوشش دستور و پوشش مسیر در این مثال وجود دارد. ممکن است خواننده با این ابهام روبه‌رو شود که در چنین حالتی، نتایج گزارش شده موردی بوده و قابل تعمیم نیست.

چنان‌چه یک دستور (بلوک پایه) جدید اجرا شود، می‌توان گفت یک مسیر جدید اجرا شده است. بنابراین از این منظر، تلاش برای افزایش پوشش بلوک پایه منجر به افزایش پوشش مسیر نیز می‌گردد. اما دو مجموعه برابر از پوشش‌های دستور، لزوماً پوشش مسیر برابری به دست نمی‌دهند. اندازه‌گیری پوشش مسیر دشوارتر بوده و سربار بیشتری به آزمون تحمیل می‌کند. از لحاظ تئوری نیز، برخی از مسیرهای ایستای برنامه، غیر قابل دسترسی هستند و با وجود حلقه‌های تکرار در برنامه، تعداد مسیرهای اجرایی در مواردی بی‌نهایت می‌شود. به سبب این‌گونه مسائل، ابزارهای معرفی شده در فصل ۲، هیچ‌کدام پوشش مسیر اجرایی را اندازه‌گیری نمی‌کنند. در بیشتر کارهای مربوط به آزمون فازی منظور از پوشش کد، همان پوشش در سطح دستورات برنامه است. بنابراین آمار و ارقام ارائه شده در این پایان‌نامه نیز مشابه کارهای پیشین برحسب همان میزان پوشش دستورات و در حالت کلی‌تر پوشش بلوک پایه بنا شده است.

در آزمایش‌هایی که پوشش دستور دو مجموعه داده آزمون، دقیقاً برابر باشد، برای مقایسه لازم است تا پوشش مسیر اندازه‌گیری گردد. در آزمایش‌های انجام شده توسط ما، همواره پوشش بلوک پایه مجموعه‌های آزمون استفاده شده پس از گرفتن اجتماع، متفاوت بوده است که نشان از متفاوت بودن پوشش مسیرها نیز دارد. لذا نتایج گزارش شده، صحت داشته و دلالت بر بهبود میزان پوشش کد SUT در حالت کلی دارد. در هر صورت، ما در نظر داریم تا آزمایش‌های خود را در مقیاس بسیار بزرگ‌تری و روی قالب‌های فایل مختلف،

تکرار کرده و اندازه‌گیری درست پوشش مسیر را نیز برای آزمایش‌های جدید انجام دهیم.

۴-۶ کارهای آتی

پایان‌نامه پیش‌رو، مباحث تولید داده آزمون، آزمون فازی و یادگیری ژرف (و به طور خاص‌تر مدل‌های زبانی عصبی) را به یکدیگر پیوند زده است. هر سه موضوع یاد شده از موضوعات مهم، کاربردی و داغ در پژوهش‌های علوم و مهندسی کامپیوتر هستند. پیشنهادهای زیادی برای کارهای آتی مرتبط با موضوع این پایان‌نامه مطرح است که در ذیل به چندین مورد از آنها اشاره می‌کنیم.

۱. استفاده از دیگر مدل‌ها و معماری‌های شبکه‌های عصبی ژرف در تولید داده آزمون. به عنوان مثال می‌توان از مدل‌های ^۱GAN [۵۳] برای تولید داده آزمون استفاده کرد. به طور خلاصه GAN از دو شبکه عصبی تشکیل شده است. یک شبکه مولد که داده‌های جدید را تولید می‌کند و یک شبکه که آنها ارزیابی می‌کند. هدف شبکه مولد این است که داده‌هایی تولید کند که دید شبکه ارزیاب خطای کمتری داشته باشند. از GAN در تولید محیط‌های جدید در بازی‌های رایانه‌ای استفاده شده است اما کاربرد آنها محدود به این مورد نیست.

۲. استفاده از یادگیری ژرف در دیگر فازرها. به عنوان نمونه یادگیری ساختار پروتکل‌های شبکه که می‌تواند برای تولید داده در فازرهای شبکه و پروتکل‌ها استفاده شود؛ این مسئله به‌ویژه در آزمون پروتکل‌هایی با ساختار ناشناخته مثل بات‌نت‌ها قابل توجه است. در این مورد یادگیری ممکن است به اهداف مهندسی معکوس کمک نماید.

۳. تولید داده آزمون براساس اهداف مختلف. هدف آزمون فازی و فازر نبایست لزوماً افزایش پوشش کد یا به عبارتی Input Gain تعیین شود. داده‌های آزمونی که توابع ناامن را فراخوانی می‌کنند برای مثال حائز اهمیت هستند [۴۹]. ایجاد مدل‌های یادگیری که امتیاز را به فراخوانی مجموعه‌ای از توابع ناامن نسبت دهند، به این امید که آزمون برنامه با ورودی‌هایی که این مسیرها را اجرا می‌کند، منجر به وقوع خطای حافظه می‌شود، می‌تواند جالب و قابل توجه باشد.

۴. افزودن حلقه بازخورد به فازر پیشنهادی در این پایان‌نامه. قبلاً نیز اشاره شد که یک محدودیت روش پیشنهادی عدم استفاده از حلقه بازخورد است. می‌توان این حلقه را در قالب یک فازر جدید اضافه

^۱Generative Adversarial Network

کرد یا برای مثال از روش پیشنهادی برای تولید دانه‌های اولیه برای فازرهای مثل AFL استفاده کرد که در این صورت بایستی یک مرحله کمینه‌سازی دانه^۱ نیز برای افزایش کارایی AFL ارایه شود.

۵. افزودن تعامل کاربر به آزمون فازی جعبه سیاه. تعامل کاربر با برنامه منجر به اجرای کدهای مربوط به کارهای کاربر می‌شود. ما در برخی آزمایش‌های خود مشاهده کردیم که این امر تأثیر زیادی بر میزان پوشش کد دارد. البته این مورد با آزمون فازی واسط کاربر متفاوت است؛ زیرا، لزوماً همه تعامل‌ها مربوط به واسط کاربر نیستند و بعضی از آنها کدهای مربوط به مرحله پرداخت فایل را اجرا می‌کنند. تعامل کاربر نیازمند دخالت مستقیم کاربر و انجام عملیاتی از طریق صفحه‌کلید یا دیگر ابزارهای ورودی است که در نتیجه خودکار نیست. بنابراین می‌توان فازری طراحی کرد که بر روی یک داده آزمون ورودی تولید شده یک دنباله از عملیات کاربر را به‌طور خودکار و به‌صورت ترتیبی یا تصادفی انجام دهد. ما این مورد را در کار بعدی خود لحاظ می‌کنیم.

۶. استفاده از یادگیری ژرف در دیگر مراحل آزمون نرم‌افزار. آزمون نرم‌افزار تنها مختص به مرحله تولید داده آزمون و مکان‌یابی خطا نیست. می‌توان در دیگر مراحل نیز از فنون یادگیری ژرف استفاده کرد. یک زمینه پژوهشی نو، ترمیم خودکار برنامه‌ها پس از مشخص شدن مکان خطا است. برای این منظور استفاده از مدل CAN^۲ [۵۴] ایده جالبی به‌نظر می‌رسد. CAN نوع خاصی از GAN است که تلاش می‌کند عوامل خلاقیت و تولید داده جدید را به آن اضافه کند. در ترمیم خودکار برنامه نیز نیازمند سازوکاری برای تنوع‌بخشی به قسمت‌های خطادار باهدف ترمیم آن قسمت‌ها، هستیم.



^۱ Seed Minimization

^۲ Creative Adversarial Networks

مراجع

- [1] A. Takanen, J. DeMott, and C. Miller. *Fuzzing for software security testing and quality assurance*. Norwood, MA, USA: Artech House, Inc., 1 ed. , 2008.
- [2] M. Sutton, A. Greene, and P. Amini. *Fuzzing: brute force vulnerability discovery*. Addison-Wesley Professional, 2007.
- [3] B. P. Miller, L. Fredriksen, and B. So, “An empirical study of the reliability of Unix utilities,” *Commun. ACM*, vol.33, pp.32–44, Dec. 1990.
- [4] B. P. Miller, D. Koski, C. Pheow, L. V. Maganty, R. Murthy, A. Natarajan, and J. Steidl, “Fuzz revisited: a re-examination of the reliability of Unix utilities and services,” tech. rep., University of Wisconsin-Madison, 1995.
- [5] J. E. Forrester and B. P. Miller, “An empirical study of the robustness of Windows NT applications using random testing,” in *Proceedings of the 4th Conference on USENIX Windows Systems Symposium - Volume 4*, WSS’00, (Berkeley, CA, USA), pp.6–6, USENIX Association, 2000.
- [6] B. P. Miller, G. Cooksey, and F. Moore, “An empirical study of the robustness of MacOS applications using random testing,” in *Proceedings of the 1st International Workshop on Random Testing*, RT ’06, (New York, NY, USA), pp.46–54, ACM, 2006.
- [7] S. Rawat, V. Jain, A. Kumar, L. Cojocar, C. Giuffrida, and H. Bos, “Vuzzer: application-aware evolutionary fuzzing,” in *Proceedings of the Network and Distributed System Security Symposium (NDSS)*, Feb. 2017.
- [8] U. Kargén and N. Shahmehri, “Turning programs against each other: high coverage fuzz-testing using binary-code mutation and dynamic slicing,” in *Proceedings of the 2015 10th Joint Meeting on Foundations of Software Engineering*, ESEC/FSE 2015, (New York, NY, USA), pp.782–792, ACM, 2015.

- [9] S. Veggiam, S. Rawat, I. Haller, and H. Bos, “Ifuzzer: an evolutionary interpreter fuzzer using genetic programming,” in *Computer Security – ESORICS 2016* (I. Askoxylakis, S. Ioannidis, S. Katsikas, and C. Meadows, eds.), (Cham), pp.581–601, Springer International Publishing, 2016.
- [10] P. Godefroid, M. Y. Levin, and D. Molnar, “Sage: whitebox fuzzing for security testing,” *Queue*, vol.10, Jan. 2012.
- [11] P. Godefroid, H. Peleg, and R. Singh, “Learn&fuzz: machine learning for input fuzzing,” in *Proceedings of the 32Nd IEEE/ACM International Conference on Automated Software Engineering*, ASE 2017, (Piscataway, NJ, USA), pp.50–59, IEEE Press, 2017.
- [12] A. Kettunen, *Test harness for web browser fuzz testing*. M.Sc. Thesis, University of Oulu, 2014.
- [13] P. Ammann and J. Offutt. *Introduction to software testing*. Cambridge University Press, 2016.
- [14] N. Rathaus and G. Evron. *Open source fuzzing tools*. Syngress Publishing, 2007.
- [15] O. Bastani, R. Sharma, A. Aiken, and P. Liang, “Synthesizing program input grammars,” *SIGPLAN Not.*, vol.52, pp.95–110, June 2017.
- [16] M. Hörschle and A. Zeller, “Mining input grammars from dynamic taints,” in *Proceedings of the 31st IEEE/ACM International Conference on Automated Software Engineering*, ASE 2016, (New York, NY, USA), pp.720–725, ACM, 2016.
- [17] I. Sutskever, O. Vinyals, and Q. V. Le, “Sequence to sequence learning with neural networks,” in *Advances in Neural Information Processing Systems 27* (Z. Ghahramani, M. Welling, C. Cortes, N. D. Lawrence, and K. Q. Weinberger, eds.), pp.3104–3112, Curran Associates, Inc., 2014.
- [18] K. Cho, B. van Merriënboer, Ç. Gülçehre, F. Bougares, H. Schwenk, and Y. Bengio, “Learning phrase representations using RNN encoder-decoder for statistical machine translation,” *CoRR*, vol.abs/1406.1078, 2014.
- [19] M. Zalewsky, “American fuzzy lop,” [Online]. Available: <http://lcamtuf.coredump.cx/afl/>, [Accessed: 2017-10-11].
- [20] M. Rajpal, W. Blum, and R. Singh, “Not all bytes are equal: neural byte sieve for fuzzing,” *CoRR*, vol.abs/1711.04596, 2017.

- [21] E. Dubrova. *Fault-tolerant design*. Springer Publishing Company, Incorporated, 2013.
- [22] Symantec, “ISTR internet security threat report,” tech. rep., 2018. <https://www.symantec.com/content/dam/symantec/docs/reports/istr-23-2018-en.pdf>.
- [23] Microsoft Corporation, “Simplified implementation of the SDL,” tech. rep., 2010.
- [24] C. Chen, B. Cui, J. Ma, R. Wu, J. Guo, and W. Liu, “A systematic review of fuzzing techniques,” *Computers and Security*, vol.75, pp.118–137, 2018.
- [25] Artifex Software, Inc., “MuPDF,” [Online]. Available: <https://mupdf.com/>, [Accessed: 2018-07-25].
- [۲۶] م. ذاکری نصرآبادی، بررسی روش‌های تولید خودکار داده‌های آزمون برای استفاده در فازهای مبتنی بر قالب فایل. سمینار کارشناسی ارشد، دانشکده کامپیوتر، دانشگاه علم و صنعت ایران، آبان ۱۳۹۶.
- [27] Paul C. Jorgensen. *Software testing a craftsman’s approach*, vol.47. CRC Press Taylor & Francis Group, fourth edi ed. , 2014.
- [28] E. Nikravan and S. Parsa, “A reasoning-based approach to dynamic domain reduction in test data generation,” *International Journal on Software Tools for Technology Transfer*, May 2018.
- [29] R. McNally, K. Yiu, and D. Grove, “Fuzzing: the state of the art,” *DSTO Defence Science and Technology Organisation*, p.55, 2012.
- [30] M. E. Khan and F. Khan, “A comparative study of white box, black box and grey box testing techniques,” *International Journal of Advanced Computer Science and Applications*, vol.3, no.6, pp.12–15, 2012.
- [31] D. L. Bruening, *Efficient, transparent, and comprehensive runtime code manipulation*. Ph.D. thesis, Cambridge, MA, USA, 2004. AAI0807735.
- [32] C.-K. Luk, R. Cohn, R. Muth, H. Patil, A. Klauser, G. Lowney, S. Wallace, V. J. Reddi, and K. Hazelwood, “Pin: building customized program analysis tools with dynamic instrumentation,” *SIGPLAN Not.*, vol.40, pp.190–200, June 2005.
- [33] QEMU Team, “QEMU,” [Online]. Available: <https://www.qemu.org/>, [Accessed: 2018-07-27].

- [34] C. Miller and Z. Peterson, “Analysis of mutation and generation-based fuzzing,” *White Paper, Independent Security Evaluators*, pp.1–7, 2007.
- [35] C. Holler, K. Herzig, and A. Zeller, “Fuzzing with code fragments,” in *Proceedings of the 21st USENIX Conference on Security Symposium, Security’12*, (Berkeley, CA, USA), pp.38–38, USENIX Association, 2012.
- [36] Microsoft, “Application verifier (appverif.exe),” [Online]. Available: <https://docs.microsoft.com/en-us/windows-hardware/drivers/debugger/application-verifier>, [Accessed: 2018-07-18].
- [37] Network Working Group, “Internet security glossary,” <https://tools.ietf.org/html/rfc2828>, [Accessed: 2017-10-11].
- [38] I. Goodfellow, Y. Bengio, and A. Courville. *Deep learning*. MIT Press, 2016. <http://www.deeplearningbook.org>.
- [39] D. Jurafsky and J. H. Martin. *Speech and language processing (3rd ed. draft)*. 2017. <https://web.stanford.edu/~jurafsky/slp3/>.
- [40] A. Karpathy, *Connecting images and natural language*. Ph.D. Thesis, Stanford University, 2016.
- [41] B. Sautermeister, *Deep learning approaches to predict future frames in videos*. M.Sc. Thesis, Technical University of Munich, 2016.
- [42] N. Srivastava, G. Hinton, A. Krizhevsky, I. Sutskever, and R. Salakhutdinov, “Dropout: a simple way to prevent neural networks from overfitting,” *Journal of Machine Learning Research*, vol.15, pp.1929–1958, 2014.
- [43] K. Greff, R. K. Srivastava, J. Koutník, B. R. Steunebrink, and J. Schmidhuber, “LSTM: a search space odyssey,” *CoRR*, vol.abs/1503.04069, 2015.
- [44] Z. C. Lipton, “A critical review of recurrent neural networks for sequence learning,” *CoRR*, vol.abs/1506.00019, 2015.
- [45] M. T. Luong, *Neural machine translation*. Ph.D Thesis, Stanford university, 2016.
- [46] T. Mikolov, *Statistical language models based on neural networks*. Ph.D. Thesis, Brno University of Technology, 2012.

- [47] M. Abadi, A. Agarwal, P. Barham, E. Brevdo, Z. Chen, C. Citro, G. S. Corrado, A. Davis, J. Dean, M. Devin, S. Ghemawat, I. J. Goodfellow, A. Harp, G. Irving, M. Isard, Y. Jia, R. Józefowicz, L. Kaiser, M. Kudlur, J. Levenberg, D. Mané, R. Monga, S. Moore, D. G. Murray, C. Olah, M. Schuster, J. Shlens, B. Steiner, I. Sutskever, K. Talwar, P. A. Tucker, V. Vanhoucke, V. Vasudevan, F. B. Viégas, O. Vinyals, P. Warden, M. Wattenberg, M. Wicke, Y. Yu, and X. Zheng, "Tensorflow: large-scale machine learning on heterogeneous distributed systems," *CoRR*, vol.abs/1603.04467, 2016.
- [۴۸] س.م. یعقوبی، طراحی و پیاده‌سازی فازر با هدف تعیین آسیب‌پذیری‌های مرورگر وب. پایان‌نامه کارشناسی ارشد، دانشکده مهندسی کامپیوتر، دانشگاه علم و صنعت ایران، آبان ۱۳۹۲.
- [۴۹] س. امینی، طراحی و پیاده‌سازی یک روش تولید داده آزمون به‌منظور کشف آسیب‌پذیری‌های نرم‌افزار. پایان‌نامه کارشناسی ارشد، دانشکده مهندسی کامپیوتر، دانشگاه علم و صنعت ایران، آبان ۱۳۹۵.
- [50] F. Chollet *et al.*, "Keras," <https://keras.io>, 2015.
- [51] D. P. Kingma and J. Ba, "Adam: a method for stochastic optimization," *CoRR*, vol.abs/1412.6980, 2014.
- [52] M. Hall, E. Frank, G. Holmes, B. Pfahringer, P. Reutemann, and I. H. Witten, "The weka data mining moftware: an update," *SIGKDD Explor. Newsl.*, vol.11, pp.10–18, Nov. 2009.
- [53] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, "Generative adversarial nets," in *Advances in Neural Information Processing Systems 27* (Z. Ghahramani, M. Welling, C. Cortes, N. D. Lawrence, and K. Q. Weinberger, eds.), pp.2672–2680, Curran Associates, Inc., 2014.
- [54] A. M. Elgammal, B. Liu, M. Elhoseiny, and M. Mazzone, "CAN: creative adversarial networks, generating "art" by learning about styles and deviating from style norms," *CoRR*, vol.abs/1706.07068, 2017.

پیوست آ

ساختار فایل PDF

در این پیوست به طور خلاصه ساختار کلی یک فایل PDF را به عنوان یک ساختار فایل پیچیده بیان می‌کنیم. در فصل ۱ اشاره کردیم که ویژگی‌های کامل قالب PDF بسیار زیاد است. بیشتر این ویژگی‌ها، در حدود ۷۰ درصد، در ارتباط با توضیح اشیای داده^۱ و ارتباط آنها بین بخش‌های مختلف یک فایل PDF هست. تمرکز اصلی در این پیوست نیز بر روی ساختار اشیای داده خواهد بود. ما ساختار کلی PDF را در دو بخش فیزیکی و منطقی بررسی خواهیم کرد.

فایل‌های PDF در یک قالب متنی کدگذاری می‌شوند که ممکن است شامل جریان دودویی^۲ مانند تصویر و غیره باشند. یک فایل PDF ترکیبی از دست‌کم یک بدنه^۳ است. یک بدنه از سه بخش تشکیل شده است: شیء^۴ (obj)، جدول ارجاع متقابل^۵ (xref) و Trailer. در ابتدای هر فایل یک سرآیند قرار می‌گیرد که با عبارت %PDF آغاز شده و در ادامه آن یک عدد که نگارش قالب PDF را مشخص می‌کند، آمده است. شکل آ-۱ یک فایل PDF شامل متن Hello Wrold را که در یک ویراستار متنی باز شده است، نشان می‌دهد. در ادامه این ساختار را بررسی می‌کنیم.

آ-۱ ساختار فیزیکی

جزئیات ساختار فیزیکی بدنه فایل PDF به شرح زیر است:

- اشیاء. داده و فراداده در پرونده‌های PDF در یک واحد اولیه که شیء نامیده می‌شود سازماندهی می‌شوند. اشیا همگی قالب مشابهی دارند که در شکل آ-۱ مشخص است و همچنین یک ساختار بیرونی مشترک هم دارند. اولین خط یک شیء شناسه آن است که برای ارجاع‌های غیر مستقیم استفاده

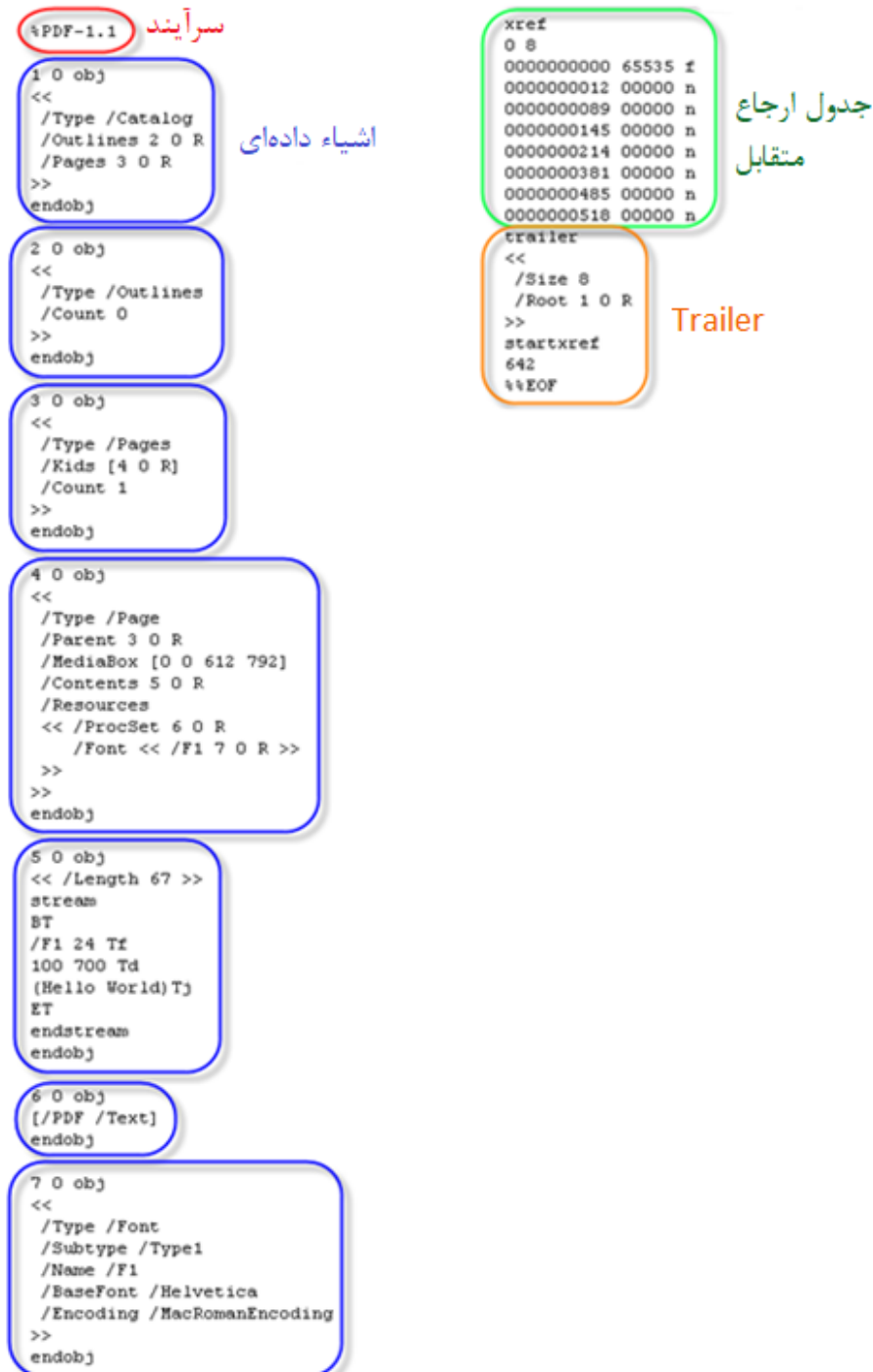
^۱Data Object

^۲Binary Stream

^۳Body

^۴Object

^۵Cross-reference Table



شکل آ-۱: یک فایل ساده PDF باز شده در یک ویراستار متنی شامل عبارت Hello Wrold. چپ: سرآیند فایل و اشیای داده‌ای؛ اشیاء با مستطیل آبی (شماره‌های ۱ تا ۷) مشخص شده‌اند. راست: ادامه محتویات همان فایل، شامل جدول ارجاع متقابل و Trailer.

می‌شود. در ادامه عدد تولیدی آن است که اگر شی با یک نسخه جدیدتر بازنویسی شود، افزایش می‌یابد. سپس رشته "obj" که شروع یک شیء را مشخص می‌کند آمده و پس از آن محتویات شیء قرار داده می‌شود. رشته "endobj" نیز پایان یافتن محدوده یک شیء را مشخص می‌کند.

● **جدول ارجاع متقابل.** این جدول در بدنه PDF شامل آدرس نسبی اشیای مورد ارجاع قرار گرفته داخل یک فایل به صورت بایت می‌باشد. در شکل آ-۱ این جدول شامل ۷ شیء با شناسه یک تا ۷ و یک مکان نگهدارنده برای شناسه صفر است که به هیچ شیئی اشاره نمی‌کند.

● **Trailer.** این قسمت از بدنه فایل PDF شامل یک فرهنگ لغت^۱ از اطلاعات بدنه، مثل تعداد کل اشیای داخل فایل (/size 8)، شناسه و شماره ترتیبی شیء ریشه (/root 10 R) و غیره است. فرهنگ لغت بین نمادهای <> و >> قرار می‌گیرد. ادامه Trailer شامل startxref است که آدرس نسبی شروع جدول ارجاع متقابل در فایل را مشخص می‌کند. این فن اجازه می‌دهد تا بدنه از پایان با خواندن startxref پویش شود، سپس به جدول ارجاع متقابل بازگشته و آن را نیز پویش می‌کند. بدین ترتیب تنها اشیای پویش می‌شوند که نیاز هستند. در شکل آ-۱ آدرس شروع جدول ارجاع متقابل که پس از startxref در Trailer آمده است برابر با ۶۴۲ است. در صورتی که آدرس‌های بایت‌های این فایل بررسی شود، مشاهده خواهد شد که بایت ۶۴۲م شروع جدول ارجاع متقابل (کاراکتر 'x' است. در پایان Trailer نیز عبارت %%EOF قرار می‌گیرد که پایان فایل را مشخص می‌کند.

ساختار جدول ارجاع متقابل و Trailer به نسبت ساده و در فایل‌های گوناگون مشابه است، اما اشیای PDF انواع مختلفی دارند. برای مثال شیء شماره ۱ در شکل آ-۱، حاوی یک ساختار فرهنگ لغت است لذا بین نمادهای <> و >> واقع شده و حاوی کلیدهایی می‌شود که با کاراکتر / شروع شده و در ادامه مقادیر آنها آمده است. مثلاً مقدار 20 R یک ارجاع به شیئی در همین فایل با شناسه ۲ و شماره ترتیبی ۰ است. از آنجایی که یک فایل ممکن است خیلی بزرگ باشد؛ آدرس نسبی شیئی که به آن ارجاع داده شده است از طریق جدول ارجاع متقابل که یک جدول دسترسی تصادفی است، قابل دستیابی خواهد بود.

اشیای PDF تنها محدود به فرهنگ لغت نمی‌شوند گرچه به نظر می‌آید که بیشتر آنها دست کم شامل یک فرهنگ لغت هستند. شیء شماره ۵ در شکل آ-۱ برای نمونه، حاوی یک جریان دودویی است که بین واژه‌های کلیدی stream و endstream قرار گرفته است. تصاویر داخل PDF نمونه‌ای از جریان‌های دودویی هستند که به این صورت می‌توانند ظاهر شوند. شکل آ-۲، یک شیء حاوی تصویر را نشان می‌دهد که در یک ویراستار متنی باز شده است. شیء شماره ۶ در شکل آ-۱ یک آرایه از کلیدها را شامل می‌شود. مقادیر آرایه می‌توانند انواع مختلفی داشته باشند که در هر صورت بین نمادهای [و] قرار می‌گیرند و با کاراکتر فاصله از یکدیگر تمیز داده می‌شوند. به دلیل تنوع ذکر شده در ساختار اشیای PDF، قواعد تعریف و ترکیب این اشیاء بیشترین بخش از توضیحات مشخصه‌های قالب فایل PDF را تشکیل می‌دهند.

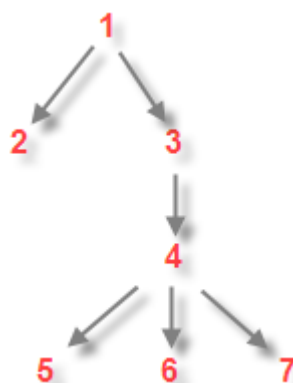
¹Dictionary

```

1397 0 obj
<</Filter/FlateDecode/Length 160>>
stream
xœ32U0P0RĐ5T02T06WH1ă*T0' EOT
SUB (CAN[ (CANéEMSTXå€Dr.-`'-34$-BIQi*-~8PDC1-34BELLP
-34S€³, !-34K'!-A,-34>34³³ ESCESC`c£İ
ÔÍ¥iœÿSš>WFFÔbgÇâéç ¶ðð\ûKÿÿêÿ;tÊİ(Ks#,,
$,ÿM>p²*z}fŽÓáNULFACKSOH.WO...@.NULÖH1£
endstream
endobj

```

شکل آ-۲: یک شی PDF حاوی یک تصویر. محتوای دودویی تصویر بین stream و endstream قرار گرفته است.



شکل آ-۳: ساختار منطقی فایل PDF نشان داده شده در شکل آ-۱. گره‌های درخت شناسه اشیاء در فایل PDF و یال‌های آن معرف نحوه فراخوانی هر شیء توسط شیء دیگر هستند.

آ-۲ ساختار منطقی

آنچه در بخش آ-۱ صحبت شد مربوط ساختار فیزیکی یک فایل PDF و نحوه قرار گرفتن اجزای فایل در کنار یکدیگر بود (مرتبط با گام پویش در هنگام اجرا). ساختار منطقی فایل PDF، یعنی قواعد حاکم بر نحوه تفسیر و پردازش آنچه در یک فایل قرار دارد و ارتباط بین اجزا، یک ساختار سلسله مراتبی است (مرتبط با گام پرداخت در هنگام اجرا). شناسه شیء ریشه در Trailer مشخص می‌شود. به همین ترتیب شیء یا اشیای بعدی مورد نیاز در شیء ریشه معلوم می‌گردد. در فایل PDF شکل آ-۱ شیء شماره ۱ ریشه است. اشیای شماره ۲ و ۳ داخل بدنه شیء ۱ مورد دسترسی قرار می‌گیرند و در نهایت ترتیب نحوه دسترسی‌ها به صورت یک درخت قابل نمایش است که ساختار منطقی فایل را نشان می‌دهد. شکل آ-۳ ساختار منطقی فایل نشان داده شده در شکل آ-۱ را نشان می‌دهد.

ساختار فیزیکی یک فایل PDF را می‌توان بدون تغییر ساختار منطقی آن، به شکل دیگری تبدیل کرد. برای مثال در شکل آ-۱ اشیاء به ترتیب صعودی شناسه خود در فایل ظاهر شده بودند. می‌شود ترتیب قرارگرفتن

```

xref
0 8
0000000000 65535 f
0000000565 00000 n
0000000509 00000 n
0000000440 00000 n
0000000273 00000 n
0000000169 00000 n
0000000136 00000 n
0000000012 00000 n

```

شکل آ-۴: جدول ارجاع متقابل بروزرسانی شده در حالتی که اشیاء به ترتیب نزولی شناسه خود در فایل PDF ظاهر شده‌اند.

اشیاء را به صورت نزولی در آورد. در این صورت جدول ارجاع متقابل بایستی بروزرسانی شود؛ زیرا، هر شی در مکان متفاوتی نسبت به قبل قرار گرفته است. اما این پدیده تأثیری بر ساختار منطقی ندارد. ساختار منطقی کماکان همان ساختار شکل آ-۳ و نتیجه اجرا نیز با فایل قبلی یکسان خواهد بود. جدول ارجاع متقابل برای حالتی که اشیاء به صورت نزولی در فایل ظاهر شده باشند، در شکل آ-۴ نشان داده شده است.

به طور کلی اشیاء PDF می‌توانند در مکان‌های تصادفی داخل یک فایل PDF ظاهر شوند بدون اینکه تأثیر بر پرداخت فایل داشته باشند. برای فایل ساده شکل آ-۱ تنها با تعویض ترتیب ظاهر شدن اشیاء می‌توان تعداد ۵۰۴۰ (برابر با ۷!) فایل با ساختار فیزیکی متفاوت داشت. این درحالی است که تغییر ترتیب قرارگیری اشیاء تنها یک روش برای جابه‌جایی ساختار فیزیکی فایل PDF است و روش‌های دیگری مثل تغییر محل قرارگیری جدول ارجاع متقابل نیز وجود دارد. بنابراین رابطه بین ساختار منطقی و فیزیکی فایل PDF به صورت یک به چند است.

آ-۳ بروزرسانی فایل PDF

دیدیم که چگونه ساختار منطقی فایل PDF مستقل از ساختار فیزیکی آن شکل می‌گیرد. بروزرسانی فایل PDF بر همین مبنا است. فایل‌های PDF می‌توانند به صورت افزایشی^۱ بروزرسانی شوند. بدین ترتیب که اگر نویسنده فایل PDF بخواهد اطلاعات داخل شیء شماره ۷ را بروزرسانی کند، یک بدنه جدید آغاز می‌کند. سپس محتوای شیء جدید را داخل آن نوشته، عدد تولیدی (عدد بعد از شناسه شیء) را یک واحد نسبت به عدد تولیدی شیء قدیم افزایش داده و برای شیء جدید می‌نویسد. در انتها نیز یک جدول ارجاع متقابل را که به شیء جدید اشاره می‌نماید، تولید کرده و بدنه جدید را به سند قبلی الصاق می‌کند.

^۱ Incremental

پیوست ب

فازر IUST DeepFuzz

در این پیوست جزئیات پیاده‌سازی و استفاده از روش پیشنهادی خود را در قالب محصول نرم‌افزاری IUST DeepFuzz شرح می‌دهیم. کلیه کدها و مستندات این محصول از طریق صفحه پروژه در وب‌سایت [GitHub](https://github.com/m-zakeri/iust_deep_fuzz)^۱ قابل مشاهده و دریافت است. برنامه به صورت مجموعه‌ای از اسکریپت‌های زبان پایتون نوشته شده است؛ لذا، روی هر سیستم عاملی پس از نصب کتابخانه‌ها و چارچوب‌های لازم قابل اجرا خواهد بود. فایل‌های پروژه، به تفکیک وظایف، در تعدادی دایرکتوری سازمان‌دهی شده است. همچنین بخشی از فازر که عملیات آزمون فازی، پایش خطا و ابزارگذاری کد SUT را انجام می‌دهد، به صورت یک فایل Batch در سیستم عامل ویندوز نوشته شده است که در صورت تغییر سیستم عامل یا هریک از ابزارهای پایش و ابزارگذاری SUT، بایستی تغییر داده شود. پیمانه تولیدکننده خودکار داده‌های آزمون اما نیاز به تغییر ندارد.

ب-۱ پیش‌نیازهای نرم‌افزاری و سخت‌افزاری

پیش‌نیازهای نرم‌افزاری پروژه در جدول **ب-۱** آمده است. پس از نصب تمامی بسته‌ها اسکریپت‌های اصلی پروژه قابل اجرا خواهد بود. برای تسریع آموزش مدل‌ها استفاده از سیستمی با پردازشگر گرافیکی قوی (تعداد هسته‌های بالا) توصیه می‌شود ولی الزامی نیست. همچنین پردازشگر مرکزی قوی و حافظه اصلی بالا (۸ گیگابایت و بیشتر) برای اجرای قابل قبول پروژه مورد نیاز خواهد بود.

ب-۲ ساختار سطح بالای پروژه

در دایرکتوری ریشه پروژه تعدادی فایل و تعدادی زیردایرکتوری قرار دارند که بخش‌های مختلف برنامه را تشکیل می‌دهند. در انتهای نام هر فایل یک عدد ترتیبی قرار دارد که نسخه آن را مشخص می‌کند. در میان فایل‌های دارای یک نام بدین ترتیب فایل با نسخه بالاتر اجرایی است و نیازی به فایل‌های قبلی نیست.

^۱https://github.com/m-zakeri/iust_deep_fuzz

جدول ب-۱: بسته‌های نرم‌افزاری مورد نیاز جهت اجرای صحیح برنامه IUST DeepFuzz.

ردیف	بسته نرم‌افزاری	نسخه مورد نیاز
۱	Python	نسخه ۳/۵ و بالاتر
۲	TensorFlow	نسخه ۱/۵ و بالاتر
۳	Keras	نسخه ۲/۲/۰

دایرکتوری‌ها عبارتند از:

- **دایرکتوری ریشه.** در این دایرکتوری ریشه فایل‌های اصلی پیکربندی^۱، فایل تعریف مدل‌های ژرف، فایل آموزش مدل و تولید داده از مدل و نیز فایل راهنما وجود دارد.
- **دایرکتوری batch_jobs.** این دایرکتوری کدهای مربوط به پیمانه‌های تزریق و پایش (خطا و پوشش کد) فازر را شامل می‌شود.
- **دایرکتوری binary_to_base64.** این دایرکتوری برای کارهای آتی رزرو شده و در نسخه فعلی محصول به‌کار نمی‌رود.
- **دایرکتوری dataset.** این دایرکتوری همان‌طور که از نام آن پیداست. محل قرار گیری مجموعه داده مورد آموزش و آزمون مدل‌ها است. برای قالب‌های فایل جدید کافی است مجموعه داده را زیر دایرکتوری مختص به آن قرار دهیم. در حال حاضر مجموعه داده خوبی برای فایل‌های PDF و مجموعه داده کوچکی نیز برای فایل‌های XML داخل آن قرار دارند.
- **دایرکتوری generated_results.** در این دایرکتوری داده‌های آزمون جدید تولید شده توسط مدل‌های ژرف ذخیره می‌شوند.
- **دایرکتوری incremental_update.** این دایرکتوری شامل کدهای مربوط به بروزرسانی افزایشی و ساخت فایل‌های PDF جدید است. بنابراین مختص آزمون فازی قالب فایل PDF است و در قالب‌های فایل دیگر کاربردی ندارد.
- **دایرکتوری logs_csv.** در این دایرکتوری فایل‌های گزارش ضبط شده در فرایند آموزش مدل‌ها ذخیره می‌شود. این فایل‌ها شامل میزان خطا، دقت و سرگشتگی مدل در هر دوره آموزش هستند.
- **دایرکتوری logs_tensorboard.** فایل‌های گزارش مربوط به ابزار مصورسازی Tensorboard در این دایرکتوری ذخیره می‌شود. این فایل‌ها سپس توسط همین ابزار خوانده و گزارش‌های دقت، خطا، معماری مدل و غیره را به‌صورت تصویری در اختیار آزمون‌گر قرار می‌دهد.

^۱ Configuration

- **دایرکتوری model_checkpoint** . در پایان هر دوره آموزش مدل، یک نمونه از مدل در این دایرکتوری ذخیره می‌شود. هر نمونه مستقل از سایر نمونه‌های بوده است. آزمون‌گر بدین ترتیب پس از اتمام فرایند آموزش می‌تواند بهترین مدل را برای تولید داده آزمون انتخاب کند.
- **دایرکتوری modelpic** . این دایرکتوری شامل گراف محاسباتی مربوط به مدل‌های ژرف است که توسط کتابخانه Keras تولید شده است.
- **دایرکتوری seed** . این دایرکتوری پوشش کد تمامی فایل‌های دانه اولیه آزمون را ذخیره می‌کند.

ب-۳ پیکربندی پروژه

در حال حاضر این محصول مبتنی بر GUI نیست و کلیه تنظیمات از طریق فایل config.py در دایرکتوری ریشه مشخص می‌شود و مهمترین فایلی است که در پروژه در اختیار آزمون‌گر قرار دارد. این فایل شامل تعدادی فرهنگ لغت است. هر فرهنگ لغت تعدادی فیلد به صورت کلید-مقدار دارد که کلید نام تنظیم ورودی بود که توسط برنامه مشخص شده است و مقدار آن توسط آزمون‌گر قبل از انجام آزمون تعیین می‌شود. فایل config.py کاملاً مستندگذاری شده و مقادیر قابل استفاده برای هر کلید در مقابل آن درج شده است. برای تغییر آن کافی است یکی از مقادیر معتبر را به جای مقدار پیش فرض قرار دهید. مهمترین تنظیم‌های قابل اعمال در این فایل عبارتند از:

- **file_fromat** . این کلید قالب فایل مورد آزمون را مشخص می‌کند. برای مثال PDF یا XML.
 - **training_set_path** . این کلید آدرس مجموعه آموزش را مشخص می‌کند.
 - **validation_set_path** . این کلید آدرس مجموعه ارزیابی را مشخص می‌کند.
 - **testing_set_path** . این کلید آدرس مجموعه آزمون را مشخص می‌کند.
 - **new_objects_pathh** . این کلید آدرس محل ذخیره داده‌های آزمون تولید شده را مشخص می‌کند.
- افزون بر تنظیمات بالا، تعدادی از تنظیمات مختص آزمون فازی قالب PDF هستند. این تنظیمات در فرهنگ لغت iu_config تعریف شده‌اند؛ از جمله:
- **single_object_update** . این کلید مقادیر True و False را می‌گیرد و همان‌طور که از نام آن پیداست تعیین می‌کند که هنگام تولید فایل PDF جدید تنها یک شیء بروزرسانی شود و یا چند شیء.
 - **portion_of_rewrite_objects** . این کلید نسبت تعداد اشیای مورد بروزرسانی را در حالت MOU مشخص می‌کند. بنابراین تنها وقتی که تنظیم single_object_update برابر False است به کار می‌رود.

- **update_policy**. این کلید نوع انتخاب اشیاء برای بروزرسانی را مشخص می‌کند که می‌تواند یکی از موارد random برای انتخاب تصادفی، bottom_up برای انتخاب برحسب شماره نزولی اشیاء و top_down برای انتخاب برحسب شماره صعودی اشیاء باشد.
- **raw_host_directory**. این کلید آدرس فایل‌های میزبان را مشخص می‌کند.

ب-۴ پیاده‌سازی‌ها

هریک از اسکریپت‌های پایتون بخشی از پروژه را پیاده‌سازی می‌کند. فایل `deep_models.py` از مدل‌های ژرف را در قالب یک تابع پایتون تعریف کرده است. فایل `lstm_text_generation_pdf_objs_9datafuzz.py` الگوریتم فاز عصبی داده و فایل `lstm_text_generation_pdf_objs_9formatfuzz.py` الگوریتم فاز عصبی فراداده را پیاده‌سازی کرده است. همچنین فایل `iu_6.py` بروزرسانی افزایشی PDF را پیاده‌سازی می‌کند.

ب-۴-۱ تعریف مدل‌ها

در Keras به دو صورت می‌توان مدل‌های یادگیری را تعریف کرد. یک روش با استفاده از پشته‌کردن لایه‌ها مختلف که برای مدل‌های ترتیبی به‌کار می‌رود و یک روش با استفاده از کلاس Model و API‌های آن، که برای ساخت مدل‌هایی با چندین ورودی و خروجی و به‌طور کلی معماری‌های پیچیده‌تر استفاده می‌شود. تعریف مدل‌های ما با روش اول انجام شده است. در برنامه **ب-۴-۱** تعریف مدل‌ها را آورده‌ایم.

```

1 from keras.models import Sequential, load_model
2 from keras.layers import Dense, Activation, Dropout
3 from keras.layers import LSTM, Bidirectional
4
5 # Unidirectional LSTM (Many to One)
6 def model_1(input_dim, output_dim):
7     print('Build model...')
8     model = Sequential()
9     model.add(LSTM(128, input_shape=input_dim))
10    model.add(Dense(output_dim))
11    model.add(Activation('softmax'))
12    return model, 'model_1'
13
```

```
14 # Unidirectional LSTM (Many to One)
15 def model_2(input_dim, output_dim):
16     model = Sequential()
17     model.add(LSTM(128, input_shape=input_dim,
18         return_sequences=True))
19     model.add(LSTM(128, input_shape=input_dim,
20         return_sequences=False))
21     model.add(Dense(output_dim))
22     model.add(Activation('softmax'))
23     return model, 'model_2'
24
25 # Unidirectional LSTM (Many to One)
26 def model_3(input_dim, output_dim):
27     model = Sequential()
28     model.add(LSTM(256, input_shape=input_dim,
29         return_sequences=True, recurrent_dropout=0.1))
30     model.add(Dropout(0.3))
31     model.add(LSTM(256, input_shape=input_dim,
32         return_sequences=False, recurrent_dropout=0.1))
33     model.add(Dropout(0.3))
34     model.add(Dense(output_dim))
35     model.add(Activation('softmax'))
36     return model, 'model_3'
37
38 # Bidirectional LSTM (Many to One)
39 def model_4_1(input_dim, output_dim):
40     model = Sequential()
41     model.add(Bidirectional(LSTM(256, return_sequences=
42         False), input_shape=input_dim, merge_mode='sum'))
43     model.add(Dense(output_dim))
```

```

39     model.add(Activation('softmax'))
40     return model, 'model_4_1'
41
42 # Bidirectional Deep LSTM (Many to One)
43 def model_4_2(input_dim, output_dim):
44     model = Sequential()
45     model.add(Bidirectional(LSTM(128, return_sequences=True),
46                             input_shape=input_dim, merge_mode='sum'))
47     model.add(Bidirectional(LSTM(128, return_sequences=False),
48                             merge_mode='sum'))
49     model.add(Dense(output_dim))
50     model.add(Activation('softmax'))
51     return model, 'model_4_2'
52 # End of model definition

```

برنامه ب-۱: پیاده‌سازی مدل‌های یادگیری در Keras

ب-۴-۲ آموزش و تولید داده

کدهای مربوط به فرایند آموزش مدل و تولید داده از مدل در یک فایل قرار دارند به ترتیب اجرا می‌شوند. در اینجا از ذکر کدها امتناع می‌کنیم؛ زیرا طولانی بوده و فایل کامل در مخزن پروژه وجود دارد. اجرای نهایی برنامه پس از انجام تنظیمات گفته شده از طریق فراخوانی هر فایل در خطفرمان سیستم عامل در مسیر پروژه انجام می‌شود. برای آموزش مدل می‌توان تعداد دوره محدودی در نظر گرفت و سپس با بررسی میزان خطا آموزش را متوقف ساخت.

ب-۴-۳ پیاده‌سازی آزمون فازی

آزمون فازی همان‌طور که گفتیم توسط یه برنامه Batch روی سیستم عامل ویندوز نوشته شده است. برنامه ب-۲ یک پیاده‌سازی اولیه از آن را نشان می‌دهد.

```
1 REM IUST DeepFuzz version 0.2
```

```
2 REM @ECHO off
```

```

3 SET PATH=%PATH%;"C:\Program Files (x86)\Microsoft Visual
   Studio 14.0\Team Tools\Performance Tools"
4 CD D:\SUT_DIRECTORY
5 CLS
6
7 START VSPerfMon /coverage /output:fuzzing_code_coverage.
   coverage
8 timeout /t 5
9
10 FOR %%i IN (.\TEST_DATA\*.pdf) DO (
11     mupdf %%i
12     timeout /t 10
13     taskkill /IM mupdf.exe
14     timeout /t 2
15 )
16
17 timeout /t 5
18 VSPerfCmd /shutdown

```

برنامه ب-۲: پیاده‌سازی آزمون فازی

ب-۵ ملاحظات عملی

در پیاده‌سازی همواره با برخی از پیچیدگی‌ها روبه‌رو می‌شویم که در تئوری مسئله وجود ندارد. در این پروژه تبدیل تمامی توالی‌های آموزشی به بردارهای عددی نیازمند حافظه بالای ۲۰ گیگابایت بود. به عبارت دیگر امکان بارگذاری یک‌باره همه مجموعه آموزش در حافظه اصلی نبود. برای حل این مشکل، هر بار یک دسته از توالی‌ها را انتخاب، آنها را از حافظه جانبی وارد حافظه اصلی کرده و پس از شرکت دادن در فرایند آموزش، این دسته را خارج و دسته بعدی را می‌خوانیم. این امکان با نوشتن از یک تابع از نوع Generator به کد برنامه اضافه گردید. مصرف حافظه به این ترتیب با تعیین اندازه دسته، در اختیار برنامه‌نویس و آزمون‌گر خواهد بود.

واژه‌نامه فارسی به انگلیسی

Representation	بازنمایی	آزمون	Test
Supervised	بانظارت	آزمون امنیت	Security Testing
Bias	بایاس	آزمون فازی	Fuzz Testing
Magic Bytes	بایت جادویی	آزمون فشار	Stress Testing
Malformed	بدشکل	آزمون قدرت‌مندی	Robustness Testing
Body	بدنه	آزمون نفوذ	Penetration Testing
Label	برچسب	آسیب‌پذیری	Vulnerability
Recognizer Program	برنامه شناسنده	آبر پارامتر	Hyper-parameter
Generator Program	برنامه مولد	ابزارگذاری	Instrumenting
Basic Block	بلوک پایه	اتکاپذیری	Dependability
Exploit	بهره‌برداری	اختلال	Noise
Overfitting	بیش‌برازش	ارزیابی بیرونی	Extrinsic Evaluation
Softmax	بیشینه هموار	ارزیابی درونی	Intrinsic Evaluation
		اشکال	Error
	پ	اعتبارسنجی	Validation
Filter	پاک‌ساز	افراز فضای ورودی	Input Space Partitioning
Monitor	پایش	افزایشی	Incremental
Render	پرداخت	افزوده	Augmented
Background Process	پردازه پس‌زمینه	الگوی آسیب‌پذیر	Vulnerable Pattern
Backpropagation	پس‌انتشار	انتهابه‌انتها	End to End
Basic Block Coverage	پوشش بلوک پایه	انفجار	Explosion
Line Coverage	پوشش خط	اینترنت چیزها	Internet of Things
Partial Line Coverage	پوشش خط جزئی		
Domain Coverage	پوشش دامنه		
Statement Coverage	پوشش دستور		
Syntax-Based Coverage	پوشش ساختار نحوی	باچ‌افزار	Ransomware

Metadata	فراداده	س
Dictionary	فرهنگ لغت	سازگار
Memory Corruption	فساد حافظه	سرگشتگی
		سروش

ق

Portability	قابلیت حمل	ش
-------------	------------	---

ک

Neural Network	شبکه عصبی	
Feed-forward Neural Network	شبکه عصبی روبه جلو	
Source Code	کد منبع	
Encoder-Decoder	کدگذار-کدگشا	
Seed Minimization	کمینه سازی دانه	
Exception Handling	کنترل استثنا	
Deep Neural Network	شبکه عصبی ژرف	
Object	شیء	
Data Object	شیء داده	

گ

Step	گام	ضد فازینگ
Time Step	گام زمانی	ضریب انشعاب
Forward Pass	گذر جلو	

ط

ل

Classification	طبقه بندی	
Design	طراحی	
Hidden Layer	لایه پنهان	

م

Generation Based	مبتنی بر تولید	عصب
Mutation Based	مبتنی بر جابه جایی	Neuron

ف

Test Set	مجموعه آزمون	فازر
Training Set	مجموعه آموزش	فازر قالب فایل
Validation Set	مجموعه ارزیابی	Fuzzer
Generative Model	مدل مولد	File Format Fuzzer

Learn&Fuzz	یادگیری و فاز	Specification	مشخصه
Unidirectional	یک‌سویه	Valid	معتبر
		Coverage Criteria	معیارهای پوشش
		Concept	مفهوم
		Regularization	منظم‌سازی
		Component	مؤلفه
		Attacker	مهاجم
		Reverse Engineering	مهندسی معکوس
		Buffer	میانگیر
		Host	میزبان

ن

Vanishing	ناپدید شدن
Learning Rate	نرخ یادگیری
Normalization	نرمال‌سازی
Markers	نشان‌گر
Requirement	نیازمندی

و

Unit	واحد
Vocabulary	واژگان
Task	وظیفه

ه

Smoothing	هموارسازی
-----------	-----------

ی

Supervised Learning	یادگیری بانظارت
Deep Learning	یادگیری ژرف

واژه‌نامه انگلیسی به فارسی

Code Coverage	پوشش کد	A	Activation Function	تابع انگیزش
Component	مؤلفه		Adaptive	تطبیق‌پذیر
Concept	مفهوم		Anti-fuzzing	ضد فازینگ
Configuration	پیکربندی		Attacker	مهاجم
Consistent	سازگار		Augmented	افزوده
Context	زمینه			
Corpus	پیکره	B	Background Process	پردازه پس‌زمینه
Cost Function	تابع هزینه		Backpropagation	پس‌انتشار
Coverage Criteria	معیارهای پوشش		Backward	روبه عقب
Cross Entropy Error	خطای آنتروپی متقاطع		Basic Block	بلوک پایه
Cross-reference Table	جدول ارجاع متقابل		Basic Block Coverage	پوشش بلوک پایه
			Bias	بایاس
D			Bidirectional	دوسویه
Data Object	شیء داده		Binary Stream	جریان دودویی
Deep Learning	یادگیری ژرف		Binary Token	توکن دودویی
Deep Neural Network	شبکه عصبی ژرف		Black Box	جعبه سیاه
Dependability	اتکاپذیری		Body	بدنه
Design	طراحی		Branch Coverage	پوشش شاخه
Dictionary	فرهنگ لغت		Branching Factor	ضریب انشعاب
Diversity	تنوع		Buffer	میانگیر
Domain Coverage	پوشش دامنه			
E		C	Encoder-Decoder	کدگذار-کدگشا
End to End	انتها به انتها		Classification	طبقه‌بندی

H

Hidden Layer لایه پنهان
History تاریخچه
Host میزبان
Hyper-parameter ابر پارامتر

I

Incremental افزایشی
Initial Seed دانه اولیه
Input Space Partitioning افراز فضای ورودی
Instrumenting ابزارگذاری
Internet of Things اینترنت چیزها
Intrinsic Evaluation ارزیابی درونی

L

Label برچسب
Learn&Fuzz یادگیری و فاز
Learning Rate نرخ یادگیری
Line Coverage پوشش خط
Logic Coverage پوشش منطق

M

Magic Bytes بایت جادویی
Malformed بدشکل
Markers نشانگر
Mean Absolute Error خطای مطلق میانگین
Mean Squared Error خطای میانگین مربعات
Memory Corruption فساد حافظه
Merge Function تابع ادغام
Metadata فراداده

Epoch دوره
Error اشکال
Error Function تابع خطا
Exception Handling کنترل استثنا
Exploit بهره‌برداری
Explosion انفجار
Extrinsic Evaluation ارزیابی بیرونی

F

Failure خرابی
Fault خطا
Feed-forward Neural Network شبکه عصبی روبه‌جلو
File Format Fuzzer فازر قالب فایل
Filter پاک‌ساز
Forward روبه‌جلو
Forward Pass گذر جلو
Fuzz Testing آزمون فازی
Fuzzer فازر

G

Generalization تعمیم‌پذیر
Generation Based مبتنی بر تولید
Generative Model مدل مولد
Generator Program برنامه مولد
Graph Coverage پوشش گراف
Gray Box جعبه خاکستری
Greedy حریصانه

Regularization	منظم‌سازی	Module	پیمانه
Render	پرداخت	Monitor	پایش
Representation	بازنمایی	Multinomial Distribution	توزیع چندجمله‌ای
Requirement	نیازمندی	Mutation Based	مبتنی بر جابه‌جایی
Reverse Engineering	مهندسی معکوس		
Robustness Testing	آزمون قدرت‌مندی		

N

S

Security Testing	آزمون امنیت
Seed Minimization	کمینه‌سازی دانه
Sequence	توالی
Shuffle	درهم‌آمیزی
Smoothing	هموارسازی
Softmax	بیشینه هموار
Source Code	کد منبع
Specification	مشخصه
Statement Coverage	پوشش دستور
Step	گام
Stress Testing	آزمون فشار
Supervised	بانظارت
Supervised Learning	یادگیری بانظارت
Syntax-Based Coverage	پوشش ساختار نحوی

T

Task	وظیفه
Test	آزمون
Test Data	داده آزمون
Test Set	مجموعه آزمون
Time Step	گام زمانی
Token	توکن
Training Set	مجموعه آموزش

Neural Network	شبکه عصبی
Neuron	عصب
Noise	اختلال
Normalization	نرمال‌سازی

O

Object	شیء
Oracle	سروش
Overfitting	بیش‌برازش

P

Parse	تجزیه
Parser	تجزیه‌گر
Partial Line Coverage	پوشش خط جزئی
Path Coverage	پوشش مسیر
Penetration Testing	آزمون نفوذ
Perplexity	سرگشتگی
Portability	قابلیت حمل
Prefix	پیشوند

R

Ransomware	باج‌افزار
Recognizer Program	برنامه شناسنده

U

Undecidable	تصمیم‌ناپذیر
Unidirectional	یک‌سویه
Unit	واحد

V

Valid	معتبر
Validation	اعتبارسنجی
Validation Set	مجموعه ارزیابی
Vanishing	ناپدید شدن
Verification	درستی‌یابی
Vocabulary	واژگان
Vulnerability	آسیب‌پذیری
Vulnerable Pattern	الگوی آسیب‌پذیر

W

Well-formed	خوش‌شکل
White Box	جعبه سفید

Abstract:

Fuzz testing (Fuzzing) is a dynamic software testing technique. In this technique with repeated generation and injection of malformed test data to the software under test (SUT), we are looking for the possible faults and vulnerabilities. To this goal, fuzz testing requires varieties of test data. The most critical challenge is to handle the complexity of the file structures as program input. Surveys have revealed that many of the generated test data in these cases follow restricted numbers and superficial paths, because of being rejected by the parser of SUT in the initial stages of parsing. Using the grammatical structure of input files to generate test data lead to increase code coverage. However, often, the grammar extraction is performed manually, which is a time consuming, costly and error-prone task. In this thesis, we proposed an automated method for hybrid test data generation. To this aim, we apply neural language models (NLMs) that are constructed by recurrent neural networks (RNNs). The proposed models by using deep learning techniques can learn the statistical structure of complex files and then generate new textual test data, based on the grammar, and binary data, based on mutations. Fuzzing the generated data is done by two newly introduced algorithms, called neural fuzz algorithms that use these models. We use our proposed method to generate test data, and then fuzz testing of MuPDF complicated software which takes portable document format (PDF) files as input. To train our generative models, we gathered a large corpus of PDF files. Our experiments demonstrate that the data generated by this method leads to an increase in the code coverage, more than 7%, compared to state of the art file format fuzzers such as American fuzzy lop (AFL). Experiments also indicate a better learning accuracy of simpler NLMS in comparison with the more complicated encoder-decoder model and confirm that our proposed models can outperform the encoder-decoder model in code coverage when fuzzing the SUT.

Keywords: Fuzz Testing, Test Data, Code Coverage, Deep Learning, Recurrent Neural Network.



**Iran University of Science and Technology
School of Computer Engineering**

Automatic Test Data Generation in File Format Fuzzers

**A Thesis Submitted in Partial Fulfillment of the Requirement for the Degree
of Master of Science in Computer Engineering**

By:

Morteza Zakeri Nasrabadi

Supervisor:

Dr. Saeed Parsa

September 2018