



سال دوم | شماره ششم | پاییز ۱۴۰۱

شاپای چاپی: ۳۶۰۷-۲۷۸۳

شاپای الکترونیکی: ۳۶۱۵-۲۷۸۳

■ صاحب امتیاز: مؤسسه آموزش عالی فردوس

مدیر مسئول: دکتر حمید طباطبایی

سرمدیر: دکتر ابراهیم محمودزاده

جانشین سرمدیر: دکتر سعیده باباجانی محمدی

مدیر داخلی: مهندس سکینه قاسمی

■ اعضای هیات تحریریه بین المللی

راجا عبدالله

استاد- گروه مهندسی کامپیوتر و سیستم های ارتباطی،

دانشکده مهندسی، دانشگاه پوترا مالزی

محمد عثمان

استاد- بخش عمومی فناوری و شبکه، دانشگاه پوترا مالزی

راجسواران لوگیس

استاد- رئیس مرکز تحلیل آسیا و اقیانوسیه، در دانشگاه

فناوری و نوآوری آسیا و اقیانوسیه.

بهمن مقیمی

استاد- دانشکده مدیریت و اقتصاد دانشگاه جرجیا در تفلیس

محمود مقومی

استاد- گروه مهندسی برق، دانشکده مهندسی، دانشگاه مالایا

مالزی

مهرداد جلالی

دانشیار- مؤسسه فناوری کارلسروهه (KIT) آلمان.

■ اعضای هیات تحریریه (به ترتیب مرتبه علمی و حروف الفبا)

پیمان اخوان

استاد- دانشگاه صنعتی قم- رئیس انجمن علمی مدیریت

دانش ایران

رضا حسنوی آتشگاه

استاد- دانشکده مهندسی صنایع دانشگاه مالک اشتر، تهران،

ایران

امیرمسعود رحمانی

استاد- دانشکده مکانیک، برق و کامپیوتر، دانشگاه آزاد

اسلامی، تهران، ایران

محمود رضایی رکن آبادی

استاد- عضو هیات امنای مؤسسه آموزش عالی فردوس و عضو

هیات علمی دانشگاه فردوسی مشهد

ابراهیم محمودزاده

استاد- دانشگاه صنعتی مالک اشتر، تهران، ایران

علی معینی

استاد- دانشکده علوم مهندسی دانشگاه تهران، ایران

محمد مهرآیین

استاد- دانشکده علوم اداری و اقتصادی دانشگاه فردوسی

مشهد، ایران

امین جاجرمی

دانشیار- گروه مهندسی برق، دانشگاه بجنورد، ایران

جواد حمیدزاده

دانشیار- دانشکده کامپیوتر و فناوری اطلاعات دانشگاه

صنعتی سجاد، مشهد، ایران

عباسعلی رضایی

دانشیار- دانشگاه پیام نور مشهد، ایران

مرتضی فرجی

دانشیار- عضو هیات امنای مؤسسه آموزش عالی فردوس و

عضو هیات علمی دانشگاه دفاع ملی تهران

محمد حسین معطر

دانشیار- دانشگاه آزاد اسلامی مشهد، ایران

سعیده باباجانی محمدی

استادیار- گروه مدیریت، مؤسسه آموزش عالی فردوس،

مشهد، ایران

علیرضا روحانی منش

استادیار- گروه مهندسی برق، دانشکده فنی و مهندسی،

دانشگاه نیشابور، ایران

محمد هادی زاهدی

استادیار- دانشگاه صنعتی خواجه نصیر طوسی، تهران، ایران

سیدکاظم شکفته

استادیار- گروه مهندسی کامپیوتر، مؤسسه آموزش عالی

شاندیدر مشهد، ایران

حمید طباطبایی

استادیار- گروه مهندسی کامپیوتر، واحد مشهد، دانشگاه آزاد

اسلامی، مشهد، ایران

مجتبی کفاشان کاخی

استادیار- گروه آموزشی علم اطلاعات و دانش شناسی دانشگاه

فردوسی مشهد، ایران

عباس مهدی زاده

استادیار- گروه کامپیوتر، مؤسسه آموزش عالی فردوس،

مشهد، ایران

ویراستار فارسی: دکتر سعیده باباجانی محمدی

ویراستار انگلیسی: دکتر عباس مهدی زاده

طراحی جلد و سرلوحه: محمدمحسن خضری

طراحی گرید و صفحه آرایی: نیما ملک زاده

کارشناس مجله: احد فانی ملکی

نشانی: ایران، مشهد، بلوار شهید کلاهدوز، شهید کلاهدوز ۳۰،

مؤسسه آموزش عالی فردوس

پایگاه اینترنتی: www.kdip.ir

تلفن: ۰۵۱-۳۷۱۳۸۰۱۱ داخلی ۷۰۳ و ۵۱-۳۷۲۹۱۱۱۴-۵

پست الکترونیکی: journal.kdip@gmail.com

مقالات مندرج لزوماً دیدگاه فصلنامه فردوس اکتشاف و پردازش هوشمند دانش نیست و مسئولیت مقالات به عهده نویسندگان است.

استفاده از مطالب و تصاویر با ذکر مأخذ بلامانع است.

پروانه انتشار فصلنامه فردوس اکتشاف و پردازش هوشمند دانش به موجب ماده ۱۳ قانون مطبوعات مصوب ۱۳۶۴/۱۲/۲۸ مجلس شورای اسلامی، از سوی اداره کل مطبوعات وزارت علوم فرهنگ و ارشاد اسلامی، طی شماره ۸۶۹۰ مورخ ۱۳۹۸/۱۰/۳۰ صادر شده است.

فهرست

سخن سردبیر

۷	زمان‌بندی وظایف در توازن منابع برای مکان‌یابی ماشین مجازی در مراکز داده ابر با استفاده از الگوریتم ژنتیک بهبود یافته
۸	مروری بر روش‌های کشف تقلب بانکی با استفاده از هوش مصنوعی
۲۰	مرکزیت در شبکه‌های اجتماعی مبتنی بر مدل آستانه خطی
۴۲	نوگرایی مصرف‌کننده و عملکرد برندهای آنلاین با توجه به نقش میانجی تعامل مشتریان با برند آنلاین
۵۴	تشخیص لبه تصاویر دیجیتالی با استفاده از خوشه‌بندی فازی بهینه شده
۶۶	واکاوی تاثیر جهت‌گیری‌های کارآفرینانه سبز بر شیوه‌های مدیریت زنجیره تامین سبز
۷۸	

دستورالعمل و راهنمای نویسندگان

● مجله «اکتشاف و پردازش هوشمند دانش» مقاله‌های منتشر نشده پژوهشی در زمینه تخصصی؛ مدیریت دانش، مدیریت فناوری، مدیریت اطلاعات را می‌پذیرد.

الف: ارسال مقاله

● جهت ارسال مقاله می‌توانید از طریق سامانه نشریه <https://www.kdip.ir> اقدام نمائید.

ب: روش نگارش

● متن مقاله بر روی فایل ساده با فرمت (A4) WORD براساس شیوه نامه فرهنگستان زبان و ادب فارسی با حروف خوانا و تیره تایپ شود. کلیه صفحات مقاله از جمله صفحاتی که شامل جداول، تصاویر و نمودارها هستند دارای قطع یکسان باشند. در متن مقاله تا حد امکان از نوشتن کلمات خارجی خودداری کلیه صفحات مقاله دارای شماره بوده و از ۲۰ صفحه تجاوز نکند.

● دقت شود که نشانه‌های نگارشی مانند؛ نقطه، ویرگول، علامت سوال، علامت تعجب و علامت نقل قول (،، ؟، !،) به کلمه قبل از خود می‌چسبند و از کلمه بعدی فاصله می‌گیرند. پرانتز، قلاب و گیومه به کلماتی که آن‌ها را در میان گرفته‌اند می‌چسبند و از کلمات قبلی یا بعدی یک فاصله دارند. فاصله بین کلمات بیش از یکی فاصله نباشد.

● برخی کلمات دارای چند جزء مختلف هستند که لازم است به صورت جدا از هم، اما در قالب یک کلمه، بیایند، مانند پیشوند، پسوند و علامت جمع (ها)، «می» مضارع و در این گونه موارد، نباید فاصله‌ای میان اجزاء کلمه باشد مثال «دست‌ها» (و نه «دست‌ها») یا «می‌شود» (و نه «می‌شود»). برای حذف فاصله بدون آن که دو حرف به هم بچسبند از کلیدهای مذکور را پشت سر هم به این ترتیب استفاده نمایید (Ctrl+ -).

ج: نحوه تهیه مقاله

● هر مقاله تخصصی بایستی تحت نرم افزار Word و دارای چکیده فارسی و لاتین با واژگان کلیدی، مقدمه، مبانی یا ادبیات موضوع و روش تحقیق، نتایج بحث، منابع مورد استفاده و یک خلاصه باشد و اصول زیر در آن رعایت شود:

۱- مشخصات نویسنده یا نویسندگان که شامل؛ نام و نام خانوادگی، سمت، محل خدمت، عنوان و درجه علمی، شماره تماس، پست الکترونیکی به فارسی و انگلیسی و تاریخ و محل انجام تحقیق می‌باشد، در یک فایل مجزا (از قسمت فایل های الحاقی یا مکمل) ارسال شود در ضمن معرفی نویسنده مسئول الزامی است.

۲- عنوان مقاله (حداکثر در ۱۲ کلمه) در وسط صفحه اول نوشته شود. اگر مقاله قسمتی از یک سری مقالات پی در پی باشد عنوان اصلی سری مقاله‌ها همراه عنوان هر قسمت و شماره ترتیب مقاله‌ها نیز ذکر گردد.

۳- چکیده در عین مختصر بودن باید محتوای مقاله را برساند. در چکیده از منابع، جداول، نمودارها و کلمات اختصاری مبهم استفاده نشود. چکیده از ۲۵۰ کلمه تجاوز نکند و تمام آن در یک پاراگراف نوشته شود.

۴- مقدمه شامل؛ اطلاعات مربوط به سابقه‌های موضوع، اهمیت تحقیق و مسأله مورد مطالعه می‌باشد.

۵- مبانی یا ادبیات موضوع، محتوای تحقیق را بر اساس منابع معتبر تبیین می‌کند.

۶- روش شناسی موضوع مورد پژوهش مشخص و روشن بیان گردد.

۷- شماره هر جدول در بالا و سمت راست آن نوشته شود. عنوان جدول گویای نتایج مندرج در آن باشد، شماره جدول در متن نیز به تناسب اشاره شود.

۸- نتایج و بحث را می‌توان به طور توأم و یا مجزا منظور کرد. بحث شامل تجزیه و تحلیل نتایج به دست آمده در ارتباط با تحقیق مورد نظر باشد.

۹- منابع مورد استفاده شامل جدیدترین اطلاعات در زمینه مورد نظر باشد. فهرست منابع به ترتیب حروف الفبایی؛ نام خانوادگی نویسندگان مقاله‌ها مرتب و شماره گذاری شود. وقتی از چند اثر مختلف یک نویسنده استفاده می‌شود ترتیب شماره گذاری این مقاله‌ها برحسب سال انتشار آنها از قدیم به جدید انجام گیرد. لازم به ذکر است کلیه منابع مورد استفاده در متن به فارسی تنظیم شده و در انتهای مقاله، ابتدا منابع فارسی به ترتیب حروف الفبایی و سپس منابع لاتین به ترتیب حروف الفبایی اشاره شود. روش منبع نویسی به صورت (APA) ای. پی. ای. باشد. لطفاً به مثال‌های زیر توجه شود.

مجلات و نشریات

نام خانوادگی، نام، (سال). عنوان مقاله، نام نشریه، (شماره جلد) شماره نشریه و صفحه‌ها.
Poh, K. W.; Yuen, P. H., & Erko, A. (2005). Entrepreneurship, innovation, and economic growth: evidence from GEM data. *Small Business Economics*, 24(3), 335-350.

کتاب

نام خانوادگی، نام، (سال انتشار). (عنوان کتاب)، (نام و نام خانوادگی مترجم)، نوبت چاپ، محل نشر: ناشر.
۱۰- چکیده انگلیسی بایستی برگردان کامل و دقیق چکیده فارسی و شامل عنوان اصلی مقاله و واژه‌های کلیدی تهیه شود.

۱۱- روش ارجاع نویسی مقالات درون‌متنی (APA) و داخل پرانتز است؛ نام خانوادگی نویسنده، سال انتشار اثر و شماره صفحه یا صفحاتی که مطلب از آن برداشته شده است، باید در متن ذکر شود (نام خانوادگی، سال، شماره صفحه). برای منابع فارسی (تألیف یا ترجمه) حتماً نام نگارنده به فارسی و سال انتشار اثر به شمسی نوشته شود و برای منابع لاتین حتماً نام به انگلیسی و سال به میلادی نوشته شود.

د: سایر موارد

۱۲- مسئولیت هر مقاله از نظر محتوای علمی و نظرات مطرح شده در متن آن، به عهده نویسنده و یا نویسندگان مسئول مقاله خواهد بود.

۱۳- تا قبل از پایان مراحل نهایی چاپ، در صورتی که مشخص گردد مقاله منتخب به هر شکلی در جای دیگری به چاپ رسیده است از انتشار آن جلوگیری خواهد شد.

۱۴- در صورتی که مقاله برای چاپ پذیرفته نشود در بخش بایگانی مجله محفوظ خواهد بود و به نویسنده برگردانده نخواهد شد.

۱۵- مقاله‌ها توسط هیأت تحریریه و با همکاری متخصصان داوری شده و در صورت تصویب بر طبق ضوابط خاص مجله به نوبت، چاپ خواهد شد.

۱۶- مجله در رد یا قبول جرح و تعدیل و ویراستاری ادبی مقاله‌ها اختیار تام دارد.

۱۷- به طور کلی به موارد زیر نیز توجه شود:

- در فایل اصلی مقاله اسم نویسنده یا نویسندگان ذکر نشود، مشخصات کامل نویسنده مسئول و نویسندگان اعم از درجه علمی، تخصص، محل کار، آدرس پستی، الکترونیکی، شماره تماس و فاکس به صورت فارسی و لاتین در یک فایل مجزا و در فایل‌های الحاقی یا مکمل ارسال شود.
- تعداد صفحات مقاله از ۲۰ صفحه بیش‌تر نباشد.
- تعداد کلمات چکیده از ۲۵۰ کلمه تجاوز ننماید.
- مقاله در صفحه ۴۴ و با تنظیمات از هر طرف ۲ سانتی متر و فاصله بین خطوط در متن مقاله ۱ باشد.

- مقاله فقط با برنامه word ۲۰۰۳ یا ۲۰۰۷، فونت متن مقاله Nazanin b سایز (اندازه) ۱۲ و فونت منابع داخل متن Times New Roman سایز (اندازه) ۱۰ و منابع پایان متن Times New Roman سایز (اندازه) ۱۱ باشد.
- عنوان مقاله به لاتین فقط کلمه اول حرف اول آن به صورت حرف بزرگ باشد و مابقی کلمات با حروف کوچک آورده شود.
- تمام اجزای مقاله در یک فایل آورده شود مانند: چکیده فارسی، لاتین، منابع، جدول ها و ...
- جدول ها و نمودارها رنگی نباشند و از کلمات و عنوان فارسی استفاده شوند.
- در منابع پایان متن از گذاشتن گیومه (« یا ») خودداری شود.
- در منابعی که سه تا پنج نویسنده دارد برای اولین بار همه نویسندگان آورده می شود و برای بار دوم از واژه «همکاران» استفاده شود.
- اگر منبعی بیش از ۶ نفر نویسنده دارد از همان ابتدا از واژه «همکاران» استفاده شود.
- در منابع داخل و پایان متن با دو نویسنده، بین نام دو نویسنده از «و در متن و» در پایان متن» استفاده شود.
- در منابع پایان متن: در منابعی که برگرفته از مقالات می باشد نام مجله به صورت ایتالیک شود. منابعی که برگرفته از کتاب می باشد نام کتاب به صورت ایتالیک شود.
- در منابعی که از نام سازمان استفاده شده، در داخل متن برای اولین بار نام کامل آن سازمان ذکر شود و برای بار دوم نام اختصاری سازمان آورده شود.
- منابع آخر متن شماره گذاری باشد (به ترتیب شماره های منابع فارسی و لاتین پشت سر هم بیاید).
- در چکیده منبع دهی مرسوم نمی باشد.
- کلید واژه فارسی بعد از چکیده فارسی قرار بگیرد و کلید واژه لاتین بعد از چکیده لاتین.
- ابتدا چکیده و واژگان فارسی، سپس چکیده و واژگان لاتین بیاید.
- منابع داخل متن آورده شود و به صورت شماره گذاری در متن نباشد در منابع داخل متن لازم نیست سال در پرانتز دیگری قرار بگیرد نام نویسنده و سال و غیره فقط در یک پرانتز قرار بگیرد.
- منابع آخر متن، شماره گذاری نباشد و اول هر منبع تورفتگی داشته باشد.
- شایان ذکر است رعایت موارد فوق، جهت قرار دادن مقاله در فرمت اولیه این نشریه بوده و به معنای پذیرش مقاله نمی باشد.
- دریافت مقاله صرفا به صورت الکترونیکی از طریق سامانه نشریه امکان پذیر است.

تنظیم خلاصه (چکیده مبسوط)

- خلاصه مقاله (در پایان مقاله) یا به عبارت دیگر، چکیده مبسوط به این ترتیب تنظیم شود.
- تعداد واژگان بکار رفته بین ۷۰۰ تا ۸۵۰ واژه باشد.
- در چکیده مبسوط، نیازی به ارایه منبع درون متن نیست.
- تیتراهای این بخش شامل موارد زیر باشد:

INTRODUCTION

THEORETICAL FRAMEWORK

METHODOLOGY

RESULTS & DISCUSSION

CONCLUSIONS & SUGGESTIONS

KEYWORDS:

REFERENCES

سخن سردبیر

«بسمه تعالی»

در موضوع الزام و ضرورت بکارگیری مدیریت دانش، در ساختارهای سازمانی و گفتمان مدیریت دانش در سازمان‌های دولتی و حتی بخش‌های خصوصی، و با عنایت به اینکه نشریه اکتشاف و پردازش هوشمند دانش، شماره ششم نشریه را منتشر نموده است، شاهد جنبه‌های مثبت بالقوه فراوانی، از منظر توجه به دارایی‌های دانشی و انتقال تجربیات و درس‌آموخته‌ها در فرآیندهای سازمانی و ارتباطات بین سازمانی هستید، که این مهم در شماره پیش رو نیز در معرض داوری اساتید و محققین قرار گرفته است. تبدیل این نکات قوت بالقوه به مزیت‌های بالفعل به بالندگی محیط‌های علمی، سازمانی و دستگاه‌های اجرایی، ارتقای بهره‌وری و توسعه نوآوری کمک می‌کند. لذا تحقق این مزیت‌ها نیازمند برنامه‌ریزی دقیق، بهره‌گیری از خرد جمعی، توانمندی‌های مراکز علمی و مشارکت کارکنان و ذی‌نفعان دستگاه‌ها است. قلب تپنده این فرایند، تبدیل از بالقوه به بالفعل، توجه و بهره‌مندی از شایستگی‌های متخصصان مدیریت و مدیریت دانش است که سالیان سال توسط دانشگاه‌ها و مراکز آموزش عالی تراز اول کشور تربیت شده و آماده به کارگیری هستند. لذا دست اندرکاران نشریه پذیرای دریافت مقالات علمی و نظرات ارزشمند محققین فرهیخته دانشگاهی هستند، باشد که با تلاش و جدیتی خالصانه، در راستای رسالت علمی که بر عهده داریم، گامی اندک و موثر برداریم. در پایان بر خود لازم می‌دانیم از تلاش‌های مستمر مدیران، اعضاء هیات علمی و دانشگاهیانی که، موسسه آموزش عالی فردوس، راپاری نمودند، داورانی که مقالات رسیده را مورد ارزیابی و بررسی قرار دادند، و آرایه دهندگان مقاله به نشریه سپاسگزاری نماییم.

مقاله پژوهشی

زمان بندی وظایف در توازن منابع برای مکان یابی ماشین مجازی در مراکز داده ابر با استفاده از الگوریتم ژنتیک بهبود یافته

Doi: 10.30508/kdip.2023.346375.1041

محسن آزاده^۱ | احمد براتیان^۲ | حمید طباطبایی (نویسنده مسئول)^۳

۱-۲-۳- گروه مهندسی کامپیوتر، واحد مشهد، دانشگاه آزاد اسلامی، مشهد، ایران.

تاریخ دریافت: ۱۴۰۱/۰۳/۱۹

تاریخ پذیرش: ۱۴۰۱/۱۰/۱۹

صفحه: ۰۰ - ۰۰

چکیده:

یکی از مهم ترین چالش های موجود در مراکز داده ابری مساله زمان بندی کارها یا اختصاص منابع به درخواست های کاربران است. دلایل متعددی ارائه شده است که این موضوع به عنوان یک مساله NP-Complete، نمود پیدا کند. در محیط محاسبات ابری هر کاربر ممکن است برای اجرای هر کار، با صدها منبع مجازی روبه رو شود. در این حال، تخصیص کارها به منابع مجازی توسط خود کاربر غیرممکن می باشد. روش های ابر اکتشافی به صورت جستجوی اکتشافی، گزینه مناسبی برای این مسائل است. از جمله این روش ها الگوریتم ژنتیک می باشد. الگوریتم ژنتیک یکی از بهترین روش های ممکن برای حل مسئله، زمان بندی وظایف در ابر می باشد که کارایی خوبی برای حل مسئله زمان بندی و تعادل بار پویا در سیستم های موازی دارد. مقاله حاضر به بررسی مفاهیم و روش های ارائه شده در این زمینه پرداخته است.

کلمات کلیدی: مکان یابی ماشین مجازی. مراکز داده ابری، زمان بندی وظایف

۱- مقدمه

عملکرد سیستم‌های کامپیوتری به چندین عامل بستگی دارد. یکی از این عوامل متوازن سازی درخواست‌ها است. این عامل به میزان کاری که در یک بازه‌ی زمانی خاص به سیستم محول شده، وابسته است. از این رو، کامپیوترها باید بر اساس اصول اولویت کارها مدیریت شوند (میشرا، ساهو و پاریدا^۱، ۲۰٪). زمان‌بندی وظایف روشی در سیستم‌های ابری است که توزیع وظایف محاسباتی در میان منابع مجازی را بر عهده دارد. زمان‌بندی وظایف موضوع مهمی است که پژوهش‌های بسیاری در این زمینه انجام شده است. با وجود اینکه ارائه‌دهندگان سرویس ابر و سرویس گیرنده‌ها، اهداف و نیازمندی‌های متفاوتی دارند ولی هدف تکنیک‌های زمان‌بندی، بهینه‌سازی یک هدف می‌باشد (دابلی، کومار و شارما^۲، ۱۸٪).

مراکز داده ابری^۳، اطلاعاتی را براساس نیاز کاربران، صرف نظر از موقعیت جغرافیایی آنها ارائه می‌دهد (ورقس و بایا^۴، ۱۸٪). مراکز داده ابری بر روی تحویل به موقع خدمات، قابلیت اطمینان، تأمین امنیت در ارسال و دریافت داده‌ها، تحمل پذیری خطا در ارائه سرویس، پایداری و ایجاد زیرساخت‌های مقیاس پذیر برای میزبانی خدمات کاربردی مبتنی بر اینترنت متمرکز شده است (لی، لیو، لین، اکسائو، زائو و کانگ^۵، ۲۱٪). عملکرد این فناوری به این صورت است که هرگونه خدماتی را با استفاده از مجازی سازی و اشتراک گذاری انجام می‌دهد. به طور کلی محاسبات ابری، مدلی است که دسترسی کاربران را بر اساس نوع تقاضاهایی که از منابع اطلاعاتی و محاسباتی دارند، امکان پذیر می‌سازد (ثیین، مائو، پروین و گانمه^۶، ۲۰٪؛ عزیزی، زند سلیمی ولی^۷، ۲۰٪). در این مدل استفاده از منابع انسانی و هزینه کاهش یافته و در عین حال، سرعت دسترسی به خدمات افزایش می‌یابد. با افزایش نرخ استفاده از این فناوری در سراسر جهان،

اختصاص یک ماشین فیزیکی به هر کاربر درخواست کننده هم از نظر هزینه و هم از نظر مصرف انرژی و تولید گرما غیرممکن است. از این رو، با گسترش محاسبات ابری، مجازی سازی به عنوان بهترین راه حل در این حوزه مورد استفاده قرار می‌گیرد. با استفاده از تکنولوژی مجازی سازی می‌توان به جای خرید چندین سرور جهت ذخیره سازی داده، داده‌های خود را بر روی مراکز داده ابری ذخیره نموده و تنها به میزان استفاده خود هزینه پرداخت کنید. در صورت نیاز، منابع بیشتری را در خدمت بگیرید (ستپاسی، آدیا، تراک، ماجهی و ساهو^۸، ۱۸٪).

کارکرد مجازی سازی به این صورت است که هر ماشین مجازی برای اجرا نیازمند یک ماشین فیزیکی است تا از طریق آن نیازهای سخت افزاری خود را (نظیر حافظه و CPU) تأمین کند. فرآیند یافتن ماشین فیزیکی با قابلیت تأمین منابع مورد نیاز برای اجرای یک ماشین مجازی را، فرآیند مکان یابی آن ماشین مجازی^۹ می‌نامیم. در محاسبات ابری، یک استخر بزرگ از منابع وجود دارد که شامل؛ ماشین‌های مجازی فراوان است. استراتژی زمان بندی، توزیع وظایف محاسباتی به استخر منابع مجازی برای به دست آوردن قدرت محاسباتی، ذخیره سازی و انواع خدمات نرم افزار با توجه به نیازهای کاربر می‌باشد.

مکان یابی ماشین‌های مجازی یکی از چالش‌های بزرگ در مراکز داده ابری است. علی‌رغم پیشرفت‌های صورت گرفته در زمینه مکان یابی ماشین‌های مجازی، هنوز مکان یابی ماشین‌های مجازی با مشکلات متعددی مواجه می‌باشد. تاکنون الگوریتم‌های مختلفی برای حل این مسئله پیشنهاد شده است (الحربی، تاین، تانگ، ژانگ، پنگ و فی^{۱۰}، ۱۹٪). در این پژوهش برای حل مسئله زمان بندی، وظایف در توازن منابع برای مکان یابی ماشین مجازی در مراکز داده ابر از الگوریتم ژنتیک بهبود یافته، بررسی می‌شود. هدف اصلی این الگوریتم‌ها برقراری

1- Mishra, Sahoo, & Parida

2- Dubey, Kumar, & Sharma

3- Cloud data centers

4- Varghese & Buyya

5- Li, Liu, Lin, Xiao, Zio, & Kang

6- Thein, Myo, Parvin, & Gawanmeh

7- Azizi, Zandsalimi, & Li

8- Satpathy, Addya, Turuk, Majhi, & Sahoo

9- Virtual Machine Placement

10- Alharbi, Tian, Tang, Zhang, Peng, & Fei

۳- مبانی نظری

یک موازنه بین هزینه انجام وظایف و زمان لازم برای انجام وظایف می باشد.

زمان بندی وظایف

زمان بندی از چالش های مهم در محاسبات ابری و سیستم های ناهمگن می باشد، به همین دلیل مطالعات و تحقیقات زیادی روی آن انجام گرفته است. در محاسبات ابری موضوع زمان بندی به عنوان یک مسئله Hard-NP مطرح می گردد (دومانال و ردی^۱، ۲۰۱۸). با استفاده از زمان بندی وظایف یک نوع موازنه بین درخواست های ورودی از سیستم برقرار می شود که سبب می شود، سیستم در شرایط پایدار و مطلوبی قرار بگیرد. در مراکز داده ابری، زمان بندی وظایف و در نتیجه تخصیص منابع، از طریق فناوری مجازی مدیریت و ارائه می شود. در محیط ابر، عملیات زمان بندی وظایف برای انتخاب و تخصیص وظایف به منابع مناسب جهت ارتقای کیفیت سرویس انجام می شود. زمان بندی وظایف به دلیل شفافیت و انعطاف پذیری سیستم ابری و نیازهای مختلف برای منابع، پیچیده تر نیز می شوند (آران آران، منجول و ساگماران^۲، ۲۰۱۹).

مکان یابی ماشین مجازی

مکان یابی ماشین مجازی به عنوان یکی از موضوعات مهم در مراکز داده ابری مطرح است و از این تکنولوژی برای مدیریت بهتر مرکز داده استفاده می نمایند. با استفاده از این تکنولوژی، ماشین های فیزیکی به چندین ماشین مجازی همراه با اشتراک گذاری منابع تبدیل می شوند در واقع مکان یابی ماشین های مجازی شامل عملیاتی برای نگاشت ماشین های مجازی بر روی ماشین های فیزیکی می باشد (لیانگ، دانگ، وانگ، و ژانگ^۳، ۲۰۲۰). در محاسبات ابری، مکان یابی ماشین های مجازی فرآیند تصمیم گیری انتخاب یک سرور برای یک ماشین مجازی، مطابق با نیاز ماشین مجازی و منابع ماشین های فیزیکی می باشد (رحمانی، خواجهوند و ترابیان^۴، ۲۰۲۰).

مکان یابی ماشین های مجازی این امکان را در

۲- اهمیت و ضرورت پژوهش

یکی از مشکلات مهم و اساسی در محیط ابری، زمان بندی وظایف می باشد. برای این منظور هر درخواست کننده براساس نیاز خود درخواست استفاده از منابع را در محیط ابر می نماید تا بتواند جریان کاری خود را توسط این منابع به انجام رساند. مشکل زمانی پیش می آید که تعداد کارها و منابع زیاد باشد، در این زمان درخواست کننده نمی تواند به راحتی منابع مناسب را انتخاب کند. این مشکل به دلیل پیچیدگی محاسباتی بالا می باشد (جانا، چاکرابرتی و ماندال^۱، ۲۰۱۹). از سوی دیگر، مراکز داده ابری جهت ارائه خدمات ابری استفاده می شوند. این مراکز برای ارائه خدمات مقادیر زیادی انرژی مصرف می کنند. علاوه بر این، مصرف سالانه انرژی را در این مراکز حدود ۹۱ میلیارد کیلووات ساعت تقریب زده شده است، که انتظار می رود این مقدار تا ۱۴۰ میلیارد کیلو وات ساعت افزایش یابد. همچنین، میزان مصرف انرژی مراکز داده ابری در سطح حدود ۱٫۱ الی ۱٫۳ درصد از کل مصرف انرژی جهان را به خود اختصاص داده، که این روند در حال افزایش است. در سال های اخیر، مصرف انرژی نه تنها هزینه برق مراکز داده را افزایش می دهد، بلکه موجب افزایش تولید و انتشار گازهای گلخانه ای و آلودگی محیط زیست شده است. از این رو، افزایش سریع در مصرف انرژی مراکز داده، هم از دیدگاه اقتصادی و هم از دیدگاه زیست محیطی تامل برانگیز است. لذا می توان با استفاده از نرم افزارهای کامپیوتری مکان یابی و تخصیص منابع را به خوبی مدیریت نمود. که برای این منظور، مقاله حاضر، به بررسی زمان بندی وظایف در توازن منابع برای مکان یابی ماشین مجازی در مراکز داده ابر با استفاده از الگوریتم ژنتیک بهبود یافته پرداخته است.

1- Jana, Chakraborty, & Mandal

2- Domanal & Reddy

3- Arunarani, Manjula, & Sugumaran

4- Liang, Dong, Wang, & Zhang,

5- Rahmani, Khajehvand, & Torabian

بقای یک فرد در نسل جدید به ترکیب خاص کروموزومی وابسته است. در مراحل تولید مثل نسل جدید ممکن است جهش‌هایی در خصوصیات یک فرد نسل جدید رخ دهد که در نتیجه فردی با ژنتیک عالی و سازگاری بالا تولید می‌شود. در روند زادوولد به گونه‌های برتر در هر نسل اجازه تولید مثل داده می‌شود و گونه‌های نامطلوب به تدریج از بین خواهند رفت و افراد نسل‌های جدید با گذشت زمان تکامل می‌یابند. مراحل الگوریتم ژنتیک به شرح زیر است (آسیما و عباس زاده اصل ۲۰۱۹، ۴):

- ایجاد جمعیت اولیه؛
- انتخاب کروموزوم‌های از جمعیت قبلی برای نسل بعدی با توجه به صحت و درستی آن و بر اساس تابع هدف برازندگی در الگوریتم؛
- تولید مثل و انجام زاد و ولد و ایجاد یک نسل جدید؛
- جهش و مشخص شدن مکان فرزند تولید شده در کروموزوم؛
- پذیرش و جا دادن فرزند جدید در داخل جمعیت؛
- جایگزینی جمعیت جدید به جای جمعیت قبلی و استفاده از جمعیت جدید در مراحل بعدی الگوریتم؛
- امتحان و در صورت دستیابی به نتایج مناسب برگشت به مرحله دوم

زمان‌بندی وظایف در توازن منابع برای مکان‌یابی ماشین مجازی در مراکز داده ابر

مکان‌یابی ماشین‌های مجازی به دلیل افزایش پیچیدگی سیستم‌های ابری اهمیت ویژه‌ای دارد. با توجه به اهمیت بالایی که فرآیند زمان‌بندی در محاسبات ابری دارد، روش‌ها و الگوریتم‌های مختلفی به منظور رسیدن به زمان‌بندی قابل قبول در محاسبات ابری ارائه شده است. در تمامی این الگوریتم‌ها سعی شده است راه حل بهینه‌ای جهت به دست آوردن یک زمان‌بندی مناسب به دست آورده شود که هر کدام از روش‌های پیاده‌سازی شده دارای معایب و مزایایی نیز بوده‌اند و یک راه حل کلی که تمام پارامترهای موجود در ابر را بهبود بخشد، پیشنهاد نشده است.

سیستم ابری فراهم می‌آورد که چندین ماشین مجازی بر روی یک میزبان فیزیکی قرار بگیرد. با توجه به افزایش روزافزون نیاز به سیستم ابری و از آنجایی که سرعت رشد داده‌ها و برنامه‌های کاربردی محاسبات ابری تا حدی زیاد است که بیش از قبل لزوم وجود سرورها و دیسک‌ها مشخص می‌شود تا این داده‌ها و برنامه‌های کاربردی را با سرعت کافی و در بازه زمانی مورد نیاز پردازش کنند، از این رو، بهینه کردن مکان‌یابی ماشین‌های مجازی همیشه یک چالش مهم محسوب می‌شود (پونراج، ۱۹، ۲؛ قره پاشا، مصدري و جعفریان، ۱۹، ۲). در نتیجه، موجب بهبود مصرف انرژی و مدیریت صحیح تخصیص منابع می‌شود. بنابراین می‌تواند نقش موثری در کاهش مصرف انرژی، جلوگیری از آلودگی محیط زیست، افزایش بهره‌وری داشته باشد.

مراکز داده ابر

محاسبات ابری یکی از روش‌ها محاسباتی است. ابرها منابع محاسباتی، بر اساس تقاضای مشتری و نیاز کاربر و بر حسب تقاضا، با دریافت هزینه از مشتری ارائه می‌دهد. ارائه دهنده خدمات ابری باید این ابر مشترک را به مشتریان خود اجاره دهد. این ابر از یک مرکز داده تشکیل شده است. به طور کلی می‌توان این گونه بیان کرد که ارائه دهنده خدمات ابری قسمتی از مرکز داده خود را به مشتریان اجاره می‌دهد. مرکز داده ابری شامل چندین سرویس دهنده است، این سرویس دهنده‌ها قدرت و منابع ابر را تشکیل می‌دهند. هر سرویس دهنده نمایانگر میزانی از منابع در مرکز داده ابری است و مجموعه‌ی سرویس دهنده‌ها با هم توانایی کلی مرکز داده را مشخص می‌کند (حسیه، لیو، بویا و زومایا^۳، ۲۰۲۰).

الگوریتم ژنتیک

الگوریتم ژنتیک نخستین بار توسط جان هلند (۱۹۶۲) ارائه شد. این الگوریتم بر اساس انتخاب طبیعی و با الهام از ژنتیک طبیعی بدن انسان رفتار می‌کند. جمعیت جدید یک جامعه از طریق زادوولد تولید می‌شوند. شانس

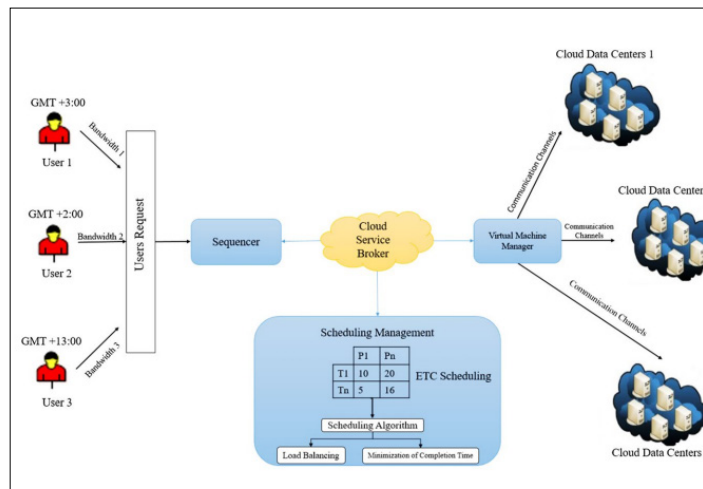
1- Ponraj

2- Gharehpasha, Masdari, & Jafarian

3- Hsieh, Liu, Buyya, & Zomaya

4- Asima, & Abbaszadeh

شكل شماره (۱)، نماي كلي زمان بندي وظيفه براي مكان يابي ماشين مجازي در محيط هاي رايانش ابري را نشان مي دهد.



شكل (۱): نماي كلي زمان بندي وظيفه در محيط هاي رايانش ابري

با استفاده الگوريتم بهينه سازي ازدحام ذرات بهبوديافته ارائه شده است. الگوريتم بهبوديافته با استفاده از يادگيري انطباقی منجر به تنوع در جمعيت می شود و لذا تعادلی بين عمليات اکتشاف و بهره برداری به دست می آید. الگوريتم زمان بندي پيشنهادهی همزمان پنج پارامتر (زمان بازگشت، بار، مصرف انرژی، هزینه و امنيت) را در حين توزيع کارها در نظر می گيرد تا در نهايت منجر به توزيع بار و کاهش مصرف انرژی می گردد. الگوريتم پيشنهادهی با استفاده از شبیه ساز کلودسيم پياده سازي و با روش هاي مربوطه (CJS, OTSS, GTSA, ISSS) مقايسه می شود. نتايج حاصل از شبیه سازي نشان می دهد، الگوريتم پيشنهادهی با در نظر گرفتن ويژگي هاي کارها و منابع، کارايي و اثربخشي قابل توجهی در محيط رايانش ابري خصوصاً در بار کاري بالا دارد.

راجوپالان، مودال و سنتي کومار^۳ (۲۰۲۰)، روش توازن بار بر مبنای الگوريتم بهينه سازي ازدحام ذرات در رايانش ابري را پيشنهاده داده اند. در اين مقاله، روشي جهت انتقال بهينه وظيفي که منجر به اضافه بار می شوند، را در انتقال

آجمال، اقبال، زيشان خان، احمد، احمد و گوپتا^۱ (۲۰۲۱) الگوريتمی ترکیبی بر پایه الگوريتم زنتیک و الگوريتم کلونی مورچگان را ارائه دادند. الگوريتم پيشنهادهی از ويژگي هاي الگوريتم زنتیک و الگوريتم کلونی مورچه ها استفاده می کند و وظيفه و ماشين هاي مجازي را به گروه هاي کوچک تر تقسيم می کند. الگوريتم پيشنهادهی به طور موثر فضای راه حل را با تقسيم وظيفه به گروه ها و با شناسايي ماشين هاي مجازي بارگذاري شده، کاهش می دهد. در اين الگوريتم با توجه به حداقل فضا، همگرایی و زمان پاسخ به طور قابل توجهی کاهش می يابد. اين روش راه حل زمان بندي عملي برای به حداقل رساندن زمان اجرای گردش کار و وظيفه ارائه می دهد، به طوري که ۶۴ درصد کاهش در زمان اجرا و ۱۱ درصد کاهش در هزینه هاي كلي مرکز داده را به همراه دارد.

منصورى، محمد حسنى زاده و غفارى^۲ (۲۰۲۱) الگوريتم زمان بندي کار مبتنی بر امنيت با استفاده از تکنیک بهينه سازي ازدحام ذرات و يادگيري انطباقی چندگانه را ارائه دادند. در اين مقاله، الگوريتم زمان بندي برای بهبود امنيت

1- Ajmal, Iqbal, Zeeshan Khan, Ahmad, Ahmad, & Gupta

2- Mansouri, Mohammad Hasani Zade, & Ghafari

3- Rajagopalan Modale, & Senthilkumar

امنیتی منابع و ارتباطات است؛ به طوری که دو محدودیت زمان و هزینه نیز برآورده شود. الگوریتم پیشنهادی -PSO WSCS که تغییراتی روی الگوریتم PSO اصلی می دهد، با سه الگوریتم دیگر زمان بندی مطرح و مشابه MPSO، VNPSO و MPSO-SA با در نظر گرفتن امنیت در محیط ابر ترکیبی مقایسه می گردد. نتایج ارزیابی حاکی از مؤثر بودن الگوریتم پیشنهادی در یافتن منابع با شباهت امنیتی نزدیک به نیازهای امنیتی می باشد. به طور متوسط، بهبود ۴۰ درصدی در نمونه های در نظر گرفته شده این مهم را نشان می دهد.

محمدزاده، مصدري، سلیمانان و جعفریان^۳ (۲۰۲۰) یک الگوریتم فرا ابتکاری بهبود یافته بر اساس الگوریتم فرا ابتکاری گرگ های خاکستری به منظور حل مسائل زمان بندی وظایف ارائه کرده اند. در الگوریتم پیشنهادی ضعیف ترین گرگ ها از جمعیت حذف شده و با گرگ های دیگری از جمعیت اولیه جاگذاری می شود. این الگوریتم با هدف بهبود عملکرد جستجو در مقابله با مسائل مختلف، افزایش سرعت هم گرایی و جلوگیری از گیر افتادن در بهینه محلی ارائه شده است. با بررسی عملکرد و مقایسه آماری نتایج به دست آمده از الگوریتم جدید با الگوریتم گرگ های خاکستری پایه و چند الگوریتم دیگر به این نتیجه رسید که با تنظیم مناسب پارامترها بهبودهای انجام شده تأثیر بسزایی در زمان بندی جریان کار در محیط محاسبات ابری دارد.

مصطفوی، احمدی و آقاسرام^۴ (۲۰۲۰) یک روش زمان بندی مبتنی بر روش یادگیری تقویتی پیشنهاد دادند که به دلیل قابلیت تطبیق با محیط و ارائه پاسخ مناسب به درخواست های متغیر با زمان، با تخصیص آینده نگرانه وظایف به منابع منجر به افزایش کارایی سیستم در بلندمدت می گردد. نتایج این مقاله نشان می دهد که روش پیشنهادی نه تنها منجر به کاهش زمان پاسخ، زمان انتظار و زمان تکمیل کار می گردد بلکه نرخ بهره وری منابع را هم به عنوان هدف فرعی افزایش می دهد. روش پیشنهادی در محدوده تعداد وظایف بالا، زمان پاسخ را به طور میانگین حدود ۴۹/۵۲ درصد نسبت به Random،

ماشین های مجازی بارگذاری شده به ماشین های مجازی در محیط رایانش ابری پیشنهاد کرده است. توابع هدف در این روش، شامل بهینه سازی اجرای وظایف و زمان انتقال در مدل بهینه سازی پیشنهادی است. آزمایش رویکرد ارائه شده در ابزارهای نرم افزار کلودسیم انجام شده است. مزیت روش ارائه شده، ارائه راه حل بهینه برای زمان بندی وظایف و توزیع برابر وظایف برای ماشین های مجازی و همچنین کاهش زمان برای فرایند تخصیص وظایف به ماشین های مجازی است.

کاشی کولایی، حسین آبادی، صائمی، شاره، سنگیاه و باین^۱ (۲۰۲۰)، روشی جهت ارتقاء زمان بندی وظایف در رایانش ابری بر مبنای الگوریتم رقابت استعماری و الگوریتم کرم شب تاب معرفی نموده اند. در روش ارائه شده، استفاده از الگوریتم هوش جمعی برای پردازش درخواست های کاربرها و زمان بندی وظایف برای توازن بار پیشنهاد شده است. هدف الگوریتم ارائه شده، بهبود زمان بندی و توازن بار در رایانش ابری با الگوریتم های جستجوی محلی و سراسری است. استفاده از الگوریتم رقابت استعماری، جهت بررسی فضای راه حل و الگوریتم کرم شب تاب، به منظور جستجوی محلی است. نتایج شبیه سازی مطابق با معیارهای پایداری، زمان پردازنده مرکزی، توازن بار، درصد راندمان و بازه زمانی نشان داده شده و الگوریتم ارائه شده قادر به تعداد وظایف بیشتری است که منجر به عملکرد بهتر الگوریتم ارائه شده نسبت به سایر الگوریتم ها شده است. این امر، به دلیل ترکیب مناسب دو الگوریتم است. مهرآوران، پژوهان و ادیب نیا^۲ (۲۰۲۰) روشی جدیدی برای زمان بندی کارها با در نظر گرفتن ملاحظات امنیتی پیشنهاد کردند. ملاحظات امنیتی شامل؛ حساسیت برای کارها که در پژوهش های اخیر در نظر گرفته شده، حساسیت برای داده های انتقالی بین کارها و همچنین ایده اصلی در نظر گرفتن قدرت امنیتی برای منابع و مسیرهای ارتباطی بین آنها می باشد. سناریوی پیشنهادی توسط الگوریتم PSO بهبود یافته (PSO-WSCS) پیاده سازی می شود. تابع هدف، حداقل کردن فاصله امنیتی کارها و داده ها از قدرت

1- Kashikolaei, Hosseinabadi, Saemi, Shareh, Sangaiah, & Bian

2- Mehravaran, Pajoohan, & Adibnia

3- Mohammadzadeh, Masdari, Soleimani, & Jafarian

4- Mostafavi, Ahmadi, & Agha Sarram

۴۶/۰۳ نسبت به Mix، ۴۳/۹۹ نسبت به FIFO، ۴۳/۵۳ نسبت به Greedy و ۳۸/۶۸ درصد نسبت به Q-sch بهبود بخشیده است.

خضرو جعفری^۱ (۲۰۱۹) بیان نمودند که روش های رایج در زمان بندی دارای معایبی از قبیل پیچیدگی زمانی بالا، هم زمان اجرا نشدن کارهای ورودی و افزایش زمان اجرای برنامه است. الگوریتم های زمان بندی بر پایه اکتشاف جهت اولویت دهی به وظایف از سیاست های متفاوتی استفاده می کنند که باعث به وجود آمدن زمان های اجرای بالا بر روی سیستم های رایانش توزیع شده ناهمگن می شود. بنابراین، روشی برای زمان بندی وظایف مناسب است که اولویت دهی آن باعث تولید زمان اجرای کل کمینه گردد. الگوریتم ژنتیک به عنوان یکی از روش های تکاملی به منظور بهینه کردن مسایل NP-Complete است. از این رو، الگوریتم ژنتیک موازی با استفاده از چارچوب نگاشت-کاهش برای زمان بندی وظایف بر روی رایانش ابری با استفاده از صف های اولویت چندگانه ارایه شده است. ایده اصلی این مقاله، استفاده از چارچوب نگاشت-کاهش برای کاهش زمان اجرای کل برنامه می باشد. بر اساس نتایج بر روی مجموعه ای از گراف های جهت دار، بدون دور تصادفی حاکی از آن است که روش پیشنهادی زمان اجرای کل دو روش موجود را با سرعت همگرایی بالا بهبود داده است. مناصراه و باعلی^۲ (۲۰۱۸) الگوریتم ترکیبی زمان بندی کار بررسی کردند. در این الگوریتم سعی بر این است که با کمک الگوریتم ژنتیک و با در نظر گرفتن تعادل بار سیستم، زمان اجرا و هزینه ای اجرا کاهش یابد. الگوریتم تلاش دارد با ایجاد تغییر در الگوریتم ژنتیک استاندارد و کاهش تعداد ایجاد جمعیت، عملکرد بهتری داشته باشد. هدف اصلی الگوریتم، تخصیص کارها به منابع با در نظر گرفتن طول کارها و تعداد دستورات قابل اجرای بوده است ماشین مجازی الگوریتم، کارها را با در نظر گرفتن تعداد کار و توانایی های منابع، به منابع تخصیص می دهد. نتایج به دست آمده نشان می دهد که الگوریتم به طور مؤثری زمان اجرا، هزینه ای اجرا و میانگین درجه عدم تعادل را بهبود بخشیده است. الگوریتم پیوندی زمان بندی کار با

دو الگوریتم بهینه سازی مورد مقایسه قرار گرفته است. در حالی که تعداد کارها افزایش پیدا می کند، مدت زمان اتمام کار نیز افزایش پیدا خواهد کرد، اما نرخ رشد در الگوریتم پیشنهادی بسیار کمتر از سایر الگوریتم ها است. دلیل این امر این است که الگوریتم پیشنهادی، زمان بندی بهینه را مبتنی بر جستجوی محلی و سراسری پیدا می کند. الگوریتم پیشنهادی کارها را با در نظر گرفتن طول کار و توانایی منبع، به منابع اختصاص می دهد. بنابراین کل زمان اجرای هر ماشین مجازی کاهش پیدا می کند.

ستاری نایینی، سالم و راشدی^۳ (۲۰۱۸) با بهره گیری از الگوریتم پرش ترکیبی قورباغه راهکاری جهت کاهش مصرف انرژی مراکز داده ابری از طریق بهینه سازی مدیریت زمان بندی کارها و ترکیب مؤثر ماشین های مجازی ارائه دادند. در این مقاله الگوریتمی جهت مدیریت زمان بندی کارها و توازن بار ارائه می شود. این الگوریتم به نام پرش ترکیبی قورباغه با بهره مندی از حافظه، همکاری و اشتراک گذاری اطلاعات بین قورباغه ها، سرعت همگرایی مناسب و انعطاف پذیری بهتر در برابر مشکل بهینه محلی، بهبود قابل توجهی نسبت به روش های سیستم تجمع مورچگان (ACO) و الگوریتم ژنتیک (GA) جهت مصرف انرژی و مهاجرت ماشین مجازی فراهم می آورد. در این مقاله، مدیریت پویای منابع بر اساس ترکیب مؤثر ماشین های مجازی انجام می شود و توسط الگوریتم پیشنهادی با توجه به قرارداد سطح سرویس پیاده سازی می شود. تفاوت این روش با روش های دیگر در این است که بهبود پارامترهای زمان، سرعت و دقت همگرایی را نشان می دهد. نتایج تجربی نشان می دهد که، روش ارائه شده موجود، از نظر مصرف انرژی، تعداد مهاجرت های ماشین مجازی و نقض قرارداد سطح سرویس عملکرد بهتری دارد.

۴- تحلیل و ارزیابی

در محاسبات ابری موضوع زمان بندی به عنوان یک مسئله Hard-NP مطرح می گردد (دومانال و همکاران، ۲۰۱۸). از این رو از اهمیت ویژه ای برخوردار است. هدف اصلی

1- Khezr, & Jafari

2- Manasrah & Ba Ali

3- Sattari Naeini, Salem, & Rashedi

کرد. الگوریتم‌های زمان‌بندی مبتنی بر توازن بار، هزینه، اولویت، کاهش زمان تکمیل کلی، مصرف انرژی، سازگاری و غیره (یزدان بخش و خرسند، ۱۳۹۹). در این بخش روش‌های موجود در زمینه زمان‌بندی وظایف و بررسی مزایا و معایب این روش‌ها دارد. در جدول شماره (۱) به بررسی و ارزیابی الگوریتم‌های معرفی شده در بخش قبل، پرداخته شده است.

۵- نتیجه‌گیری

زمان‌بندی وظیفه یک نقش کلیدی در سیستم‌های

از زمان‌بندی وظایف، اولویت‌بندی و تخصیص وظایف به منابع آزاد است به گونه‌ای که حداکثر وظایف به صورت موازی در حال پردازش باشند. دیگر اهداف زمان‌بندی شامل: تعادل بار، کیفیت خدمات، اصل اقتصادی بودن، بهترین زمان اجرا و توان عملیاتی سیستم است (ابولقاه و دیابات، ۲۰۲۱). روش‌ها و طریق‌های زمان‌بندی متفاوتی برای تخصیص وظایف به گره‌های مختلف و ماشین‌های مجازی در دسترس در محیط‌های ابری به استخدام گرفته می‌شوند. الگوریتم‌های زمان‌بندی ارائه شده را می‌توان مبتنی بر پارامترهای بهبود در محیط ابری دسته‌بندی

جدول (۱): مقایسه الگوریتم‌های ارائه شده			
نویسندگان	عنوان مقاله	مزایا	معایب
آجمال و همکاران (۲۰۲۱)	الگوریتم ترکیبی ژنتیک مورچگان برای زمان‌بندی کارآمد در مراکز داده ابری	۱. استفاده از ویژگی‌های الگوریتم ژنتیک و الگوریتم کلونی مورچه‌ها ۲. کاهش فضای راه حل با تقسیم وظایف ۳. کاهش زمان پاسخ‌گویی ۴. سرعت همگرایی بالا ۵. کاهش زمان اجرا ۶. کاهش در هزینه‌های کلی مرکز داده ۷. کاهش توان محاسباتی	۱. عدم توجه به مصرف انرژی ۲. عدم توجه به استفاده از منابع ۳. عدم توجه به مقیاس‌بندی مرکز داده ۴. هم زمان اجرا نشدن کارها
منصوری و همکاران (۲۰۲۱)	الگوریتم زمان‌بندی کار مبتنی بر امنیت با استفاده از تکنیک بهینه‌سازی ازدحام ذرات و یادگیری انطباقی چندگانه	۱. کاهش زمان پاسخ‌گویی ۲. توزیع بار بهینه ۳. کاهش مصرف انرژی	۱. عدم توجه به مقیاس‌بندی مرکز داده ۲. عدم توجه به کاهش فضای راه حل
راجوپلان و همکاران (۲۰۲۰)	زمان‌بندی بهینه وظایف در رایانش ابری با استفاده از الگوریتم ترکیبی کرم شب تاب-ژنتیک	۱. ارائه راه‌حل بهینه برای زمان‌بندی وظایف ۲. توزیع برابر وظایف برای ماشین‌های مجازی ۳. کاهش زمان برای فرایند تخصیص وظایف به ماشین‌های مجازی ۴. قابلیت جستجوی مؤثر و سریعی در یک محیط ابری پویا ۵. سرعت همگرایی بالا ۶. تخصیص بهینه منابع ۷. کاهش زمان اجرا	۱. عدم توجه به کاهش فضای راه حل ۲. عدم توجه به مقیاس‌بندی مرکز داده ۳. عدم توجه به توان محاسباتی ۴. هم زمان اجرا نشدن کارها

جدول (۱): مقایسه الگوریتم‌های ارائه شده

نویسندگان	عنوان مقاله	مزایا	معایب
کاشی کولایی و همکاران (۲۰۲۰)	بهبود زمان بندی وظایف در رایانش ابری بر اساس الگوریتم رقابتی امپریالیستی و الگوریتم فرقی	۱. بهبود زمانبندی ۲. توازن بار در رایانش ابری ۳. کاهش زمان مصرفی CPU ۴. افزایش تعادل بار ۵. افزایش درصد کارایی و ماندگاری	۱. عدم توجه به کاهش فضای راه حل ۲. عدم توجه به مقیاس بندی مرکز داده ۳. عدم توجه به توان محاسباتی ۴. عدم توجه به سرعت همگرایی ۵. زمان بالای تخصیص وظایف به ماشین‌های مجازی ۶. هم زمان اجرا نشدن کارها
مهرآوران و همکاران (۲۰۲۰)	زمان بندی گردش کار در محیط ابر ترکیبی با در نظر گرفتن امنیت کارها و ارتباطات با الگوریتم ازدحام ذرات بهبود یافته	۱. کاهش زمان اجرا ۲. کاهش در هزینه‌های کلی مرکز داده ۳. حداقل کردن فاصله امنیتی کارها و داده‌ها	۱. عدم توجه به کاهش فضای راه حل ۲. عدم توجه به مقیاس بندی مرکز داده ۳. عدم توجه به سرعت همگرایی
محمدزاده و همکاران (۲۰۲۰)	ارائه یک الگوریتم بهبود یافته بهینه سازی گرگ‌های خاکستری برای زمان بندی جریان کار در محیط محاسبات ابری	۱. کاهش زمان اجرا ۲. افزایش سرعت همگرایی ۳. جلوگیری از افتادن در بهینه محلی	۱. عدم توجه به مقیاس بندی مرکز داده ۲. عدم توجه به کاهش فضای راه حل ۳. عدم توجه به مصرف انرژی
مصطفوی و همکاران (۲۰۲۰)	زمان بندی آینده‌نگرانه وظایف در محاسبات ابری مبتنی بر یادگیری تقویتی	۱. کاهش زمان پاسخ ۲. کاهش زمان انتظار ۳. کاهش زمان تکمیل کار ۴. افزایش نرخ بهره‌وری منابع	۱. عدم توجه به مقیاس بندی مرکز داده ۲. عدم توجه به کاهش فضای راه حل ۳. عدم توجه به مصرف انرژی
مناصراه و همکاران (۲۰۱۸)	زمان بندی گردش کار با استفاده از الگوریتم ترکیبی GA-PSO در رایانش ابری	۱. کاهش زمان اجرا ۲. کاهش هزینه اجرا ۳. تخصیص کارها به منابع با در نظر گرفتن طول کارها و توانایی‌های منابع ۴. بهبود میانگین درجه عدم تعادل ۵. توجه به مقیاس بندی مرکز داده	۱. هم زمان اجرا نشدن کارها ۲. عدم توجه به سرعت همگرایی ۳. فضای بالای راه حل
خضرو جعفری (۲۰۱۹)	زمان بندی کارها در محیط‌های ابری با استفاده از چارچوب نگاشت - کاهش و الگوریتم ژنتیک	۱. هم زمان اجرا شدن کارها ۲. کاهش پیچیدگی زمانی ۳. کاهش زمان اجرای برنامه ۴. سرعت همگرایی بالا	۱. عدم توجه به مقیاس بندی مرکز داده ۲. عدم توجه به کاهش فضای راه حل ۳. عدم توجه به استفاده از منابع
ستاری و همکاران (۲۰۲۰)	بهره‌گیری از الگوریتم پرش ترکیبی قورباغه جهت کاهش مصرف انرژی مراکز داده ابری از طریق بهینه‌سازی مدیریت زمان بندی کارها و ترکیب مؤثر ماشین‌های مجازی	۱. کاهش زمان اجرا ۲. سرعت همگرایی بالا ۳. انعطاف پذیری بالا در برابر مشکل بهینه محلی ۴. کاهش مصرف انرژی ۵. تعداد بهینه مهاجرت‌های ماشین مجازی	۱. عدم توجه به کاهش فضای راه حل ۲. عدم توجه به مقیاس بندی مرکز داده

رسیدگی کردن و انجام آنها، فرآیند زمان بندی و تخصیص منابع کاری دشوار می شود و باید از طریق الگوریتم مناسب این زمان بندی بین فراهم کننده منابع و درخواست کاربران انجام شود. در اینجا روش متفاوت را برای این نوع زمان بندی وظایف با الگوریتم ژنتیک بررسی گردید. که هر کدام از این روش ها دارای نقاط ضعف و قوت خاص خودشان هستند. الگوریتم های زمان بندی موجود در این پژوهش اهدافی دارند که عبارتند از: به حداقل رساندن کل زمان اجرای کار، افزایش قدرت پردازش، حداقل رساندن کل زمان پردازش، به حداقل رساندن طول برنامه، کاهش مصرف انرژی، دستیابی توان بالای سیستم و غیره است. زمان بندی در سیستم های رایانش توزیع شده یک مسئله مهم می باشد و الگوریتم های زمان بندی گوناگونی در این زمینه ارائه شده است. پیشنهاد می شود به عنوان پژوهش های بعدی از سایر الگوریتم های فرابتکاری به منظور زمان بندی وظایف در مراکز داده های ابری استفاده گردد و نتایج حاصله با الگوریتم ژنتیک ارائه در این پژوهش مقایسه گردد. همچنین، می توان جهت کاهش پیچیدگی زمانی محاسباتی الگوریتم ژنتیک ارائه شده از روش موازی سازی استفاده نمود. پردازش موازی یکی از مهم ترین برتری های الگوریتم ژنتیک می باشد، به این معنی که در این روش به جای یک متغیر، در یک زمان یک جمعیت را به سوی نقطه بهینه رشد می دهید. بنابراین سرعت همگرایی روش بسیار بالا می رود.

محاسباتی ابر بازی می کند و این نوع زمان بندی بر اساس یک معیار تنها نمی تواند انجام شود، بلکه تعداد زیادی قوانین و شرایط به صورت یک توافق بین کاربران و فراهم کنندگان ابر باید در نظر گرفته شوند. در واقع این توافق، کیفیت سرویسی است که کاربران از فراهم کنندگان انتظار دارند. مراکز داده ابر نه تنها باید این وظیفه های عظیم را اجرا کنند بلکه باید نیازمندی های چندگانه کاربران مختلف را ارضا کنند. به عنوان نمونه در مقاله منصوری و همکاران (۲۰۲۱) و مقاله مهرآوران و همکاران (۲۰۲۰) عنصر امنیت در زمان بندی وظایف در نظر گرفته شده است در حالی که در سایر مقالات به این عنصر، توجهی نشده است. یا در الگوریتم ترکیبی بر پایه الگوریتم ژنتیک و الگوریتم کلونی مورچگان که توسط آجمال و همکاران (۲۰۲۱) ارائه شده به طور موثر فضای راه حل را با تقسیم وظایف به گروه ها و با شناسایی ماشین های مجازی بارگذاری شده کاهش می دهد. همچنین، بررسی نتایج مقالات نشان می دهد که الگوریتم های ترکیبی کارایی و کیفیت بالاتری را نتیجه می دهند. از سوی دیگر الگوریتم پرش ترکیبی قورباغه ارائه شده توسط ستاری و همکاران (۲۰۱۸)، سرعت همگرایی مناسب و انعطاف پذیری بهتر در برابر مشکل بهینه محلی، بهبود قابل توجهی نسبت به روش های سیستم تجمع مورچگان (ACO) و الگوریتم ژنتیک (GA) جهت مصرف انرژی و مهاجرت ماشین مجازی فراهم می آورد. زمانی که وظایف و درخواست ها برای منابع زیاد شود،

منابع

- 6-Abualigah, L., Diabat, A. (2021). A novel hybrid antlion optimization algorithm for multi-objective task scheduling problems in cloud computing environments. *Cluster Comput*, 24(): 205-223.
- 7-Ajmal, M.S., Iqbal, Z., Zeeshan Khan, F., Ahmad, M., Ahmad, I., Gupta, B.B. (2021). Hybrid ant genetic algorithm for efficient task scheduling in cloud data centers, *Computers & Electrical Engineering*, 95(): 107-119.
- 8-Alharbi, F., Tian, Y.C., Tang, M., Zhang, w.Z., Peng, C., Fei, M. (2019). An Ant Colony System for energy-efficient dynamic Virtual Machine Placement in data centers, *Expert Systems with Applications*, 120(): 228-238.
- 9-Arunarani, A.R., Manjula, D., Sugumaran, V. (2019). Task scheduling techniques in cloud computing:

- A literature survey, *Future Generation Computer Systems*, 91(): 407-415.
- 10-Asima, M., Ali Abbaszadeh Asl, A. (2019). Developing a Hybrid Model to Estimate Expected Return Based on Genetic Algorithm. *Financial Research Journal*, 21(1), 101-120.
- 11-Azizi, S., Zandsalimi, M., Li, D. (2020). An energy-efficient algorithm for virtual machine placement optimization in cloud data centers. *Cluster Comput* 23(): 3421-3434.
- 12-Domanal, S.G., Reddy, G.R.M. (2018). An efficient cost optimized scheduling for spot instances in heterogeneous cloud environment. *Future Generation Computer Systems*, 84(): 11-21.
- 13-Dubey, K., Kumar, M., Sharma, S. C. (2018). Modified HEFT Algorithm for Task Scheduling in Cloud Environment, *Procedia Comput. Sci.*, 125(), 725-732.
- 14-Gharehpasha, S., Masdari, M., Jafarian, A. (2019). A New Approach for Optimal Placement of Virtual Machines in Cloud Datacenters Using Discrete Gravitational Search Algorithm and Chaotic Functions. *Journal of Soft Computing and Information Technology*, 8(2): 44-54.
- 15-Gharehpasha, S., Masdari, M., Jafarian, A. (2020). The Placement of Virtual Machines Under Optimal Conditions in Cloud Datacenter. *Information Technology and Control*, 48(4).
- 16-Hsieh, S.H., Liu, C.S., Buyya, R., Zomaya, A.Y. (2020). Utilization-prediction-aware virtual machine consolidation approach for energy-efficient cloud data centers, *Journal of Parallel and Distributed Computing*, 139(): 99-109.
- 17-Jana, B., Chakraborty, M., Mandal, T. (2019). A Task Scheduling Technique Based on Particle Swarm Optimization Algorithm in Cloud Environment. In: Ray K., Sharma T., Rawat S., Saini R., Bandyopadhyay A. (eds) *Soft Computing: Theories and Applications. Advances in Intelligent Systems and Computing*, Springer, Singapore, 742(): 525-536.
- 18-Jiang, Y., Wu, P., Zeng, J., Zhang, Y., Zhang, Y., Wang, S. (2019). Multiparameter and multi-objective optimisation of articulated monorail vehicle system dynamics using genetic algorithm, *Veh. Syst. Dyn*: 1-18.
- 19-Kashikolaei, S.M.G., Hosseinabadi, A.A.R., Saemi, B., Shareh, M.B., Sangaiah, A.K. and Bian, G.B. (2020). An enhancement of task scheduling in cloud computing based on imperialist competitive algorithm and firefly algorithm. *The Journal of Supercomputing*, 76(8): 6302-6329.
- 20-Khezr, N., Jafari Novimipour, N. (2019). Scheduling tasks in cloud environments using mapping framework - reduction and genetic algorithm Research Areas: General, *Journal of Information and Communication Technology*, 10(37): 71-84.
- 21-Li, X.Y., Liu, Y., Lin, Y.H., Xiao, L.H., Zio, E., Kang, R. (2021). A generalized petri net-based modeling framework for service reliability evaluation and management of cloud data centers, *Reliability Engineering & System Safety*, 207(): 107-181.
- 22-Liang, B., Dong, X., Wang, Y., Zhang, X. (2020). Memory-aware resource management algorithm for low-energy cloud data centers, *Future Generation Computer Systems*, 113(): 329-342.
- 23-Manasrah, A., Ba Ali, H. (2018). Workflow Scheduling Using Hybrid GA-PSO Algorithm in Cloud Computing *Hindawi January 2018 Wireless Communications and Mobile Computing*, 8(3):1-16 .
- 24-Mansouri, N., Mohammad Hasani Zade, B., Ghafari, R. (2021). Security-aware Task Scheduling Algorithm based on Multi Adaptive Learning and PSO Technique. *Electronic and Cyber Defense*, 9(2), 159-178.
- 25-Mehravar, M., Pajoohan, M., Adibnia, F. (2020). Secure and confidential workflow scheduling in hybrid cloud using improved particle swarm optimization algorithm. *Electronic and Cyber Defense*, 7(4): 131-145.
- 26-Mishra, S.K., Sahoo, B., Parida, P.P. (2020). Load balancing in cloud computing: A big picture, *Journal of King Saud University - Computer and Information Sciences*, 32(2): 149-158.
- 27-Mohammadzadeh, A., Masdari, M., Soleimani Gharehchopogh, F., Jafarian, A. (2020). An Improved Grey Wolves Optimization Algorithm For Workflow Scheduling In Cloud Computing Environment. *Journal of Soft Computing and Information Technology*, 8(4): 17-29.
- 28-Mostafavi, S., Ahmadi, F., Agha Sarraam, M. (2020). Reinforcement-Learning-based Foresighted

- Task Scheduling in Cloud Computing. *Tabriz Journal Of Electrical Engineering*, 50(1), 387-401.
- 29-Ponraj, A. (2019). Optimistic virtual machine placement in cloud data centers using queuing approach, *Future Generation Computer Systems*, 93(): 338-344.
- 30-Rahmani, S., Khajehvand, V., Torabian, M. (2020). Burst-aware Placement for Improving VM Consolidation in Cloud Environment. *Journal of Soft Computing and Information Technology*, 9(3): 1-14.
- 31-Rajagopalan, A., Modale, D.R., Senthilkumar, R. (2020). Optimal scheduling of tasks in cloud computing using hybrid firefly-genetic algorithm. In *Advances in decision sciences, image processing, security and computer vision*. Springer, Cham: 678-687.
- 32-Satpathy, A., Addya, S.K., Turuk, A.K., Majhi, B., Sahoo, G. (2018). Crow search based virtual machine placement strategy in cloud data centers with live migration, *Computers & Electrical Engineering*, 69(): 334-350.
- 33-Sattari Naeini, V., Salem, Y., Rashedi, E. (2018). Application of Shuffled Frog-Leaping Algorithm to Reduce Energy Consumption in Cloud Data Centers by Optimizing Scheduling Management and Virtual Machines Consolidation. *Tabriz Journal Of Electrical Engineering*, 48(2), 687-698.
- 34-Thein, T., Myo, M.M., Parvin, S., Gawanmeh, A. (2020). Reinforcement learning based methodology for energy-efficient resource allocation in cloud data centers, *Journal of King Saud University - Computer and Information Sciences*, 32(10): 1127-1139.
- 35-Varghese, B., Buyya, R. (2018). Next generation cloud computing: New trends and research directions, *Future Generation Computer Systems*, 79(): 849-861.
- 36-Yazdanbakhsh, M., Khorsand, R. (2019). A Task Scheduling Strategy to Improve Qualitative Features in the Cloud Computing Environment. *Tabriz Journal Of Electrical Engineering*, 49(3), 1427-1437.

©Authors, Published by Journal of Intelligent Knowledge Exploration and Processing. This is an open-access paper distributed under the CC BY (license <http://creativecommons.org/licenses/by/4.0/>).



مقاله پژوهشی

مروری بر روش‌های کشف تقلب بانکی با استفاده از هوش مصنوعی

Doi: 10.30508/kdip.2022.349333.1044

سیدمجید اکبرالسادات (نویسنده مسئول)^۱ | بابک اسماعیل پور^۲

۱- گروه هوش مصنوعی و رباتیک، دانشگاه آزاد اسلامی (واحد قوچان)، قوچان، ایران

۲- استادیار گروه کامپیوتر، دانشگاه آزاد اسلامی (واحد مشهد)، مشهد، ایران

تاریخ دریافت: ۱۴۰۱/۰۴/۰۶

تاریخ پذیرش: ۱۴۰۱/۱۰/۰۶

صفحه: ۵۵ - ۵۰

چکیده:

تقلب بانکی با توجه به تأثیرات مخرب آن، یکی از تهدیدآمیزترین مشکلاتی است که هر جامعه بشری با آن دست‌وپنجه نرم می‌کند. این عمل به استفاده عمدی از اطلاعات نادرست برای کلاهبرداری از پول یا دارایی فرد یا سازمان دیگری اشاره دارد. صنعت بانکداری برای چندین دهه از سیستم‌های مبتنی بر قوانین برای شناسایی تقلب و بررسی انسانی تراکنش‌ها استفاده کرده است. سیستم‌های مبتنی بر قانون شامل الگوریتم‌هایی هستند که انواع مختلفی از اقدامات شناسایی را انجام می‌دهند که به صورت دستی توسط متخصصان تقلب نوشته شده‌اند. این سیستم‌ها به تنظیم دستی سناریوها نیاز دارند، که در تشخیص ضمنی همبستگی‌های معاملاتی که به تقلب اشاره می‌کنند، دچار چالش می‌شوند. با توجه به ضعف‌های ذاتی رویکرد تشخیص تقلب مبتنی بر قانون در بانک‌ها و داده‌های محدودی که بر الگوریتم‌های یادگیری ماشین تحت نظارت متداول استفاده می‌شود، نیاز مبرمی به تکنیک‌ها یا سیستم‌های جدید تشخیص تقلب وجود دارد که بتوانند با افزایش سریع تقلب‌ها و موارد پول‌شویی مقابله کنند. این تحقیق با استفاده از رویکرد مروری و با هدف بررسی روش‌های کشف تقلب بانکی به بررسی ادبیات تحقیق و توصیف نتایج مرتبط می‌پردازد.

کلمات کلیدی: تقلب بانکی، بانکداری الکترونیک، کلاهبرداری، معاملات متقلبانه، تراکنش‌های جعلی، کشف تقلب بانکی.

۱- مقدمه

تقلب با توجه به تأثیرات مخرب آن، یکی از تهدیدآمیزترین مشکلاتی است که هر جامعه بشری با آن دست‌وپنجه نرم می‌کند. این عمل به استفاده عمدی از اطلاعات نادرست برای کلاهبرداری از پول یا دارایی یک فرد یا سازمان دیگر اشاره دارد (ویلیام^۱، ۲۰۲۲). در طول رکود اقتصادی ناشی از همه‌گیری کرونا، به دلیل افزایش تعداد افراد بیکار در طول دوره، شاهد افزایش فعالیت‌های کلاهبرداری بوده‌ایم (شیهم بیتسا^۲، ۲۰۲۱). هزاران نفر بیکار شدند، حقوق‌ها کاهش یافت و نرخ بیکاری افزایش یافت. به طور طبیعی، افراد بیشتری منابع کمتری برای بقای خود در اختیار دارند، که این انگیزه‌ای برای افزایش فعالیت‌های کلاهبرداری می‌باشد، زیرا مردم در تلاش برای زنده ماندن در دوران سخت اقتصادی هستند. کارشناسان تقلب استدلال می‌کنند که تقلب از سه عنصر که شامل فشار، فرصت و منطقی‌سازی است؛ نشأت می‌گیرد (کارتیس^۳، ۲۰۱۲). به گفته محققین، کلاهبرداری زمانی اتفاق می‌افتد که فردی فشار و انگیزه غیرقابل تقسیم برای ارتکاب کلاهبرداری ایجاد کند. در بیشتر موارد، برای مرتکب تقلب، نیازهای برآورده نشده بی‌پایان است و برای افراد مختلف متفاوت می‌باشد. که می‌تواند شامل؛ افزایش صورت‌حساب پزشکی، کاهش درآمد در خانواده یا بدهی قمار باشد. هنگامی که فرد نیازهای برآورده نشده داشته باشد و منابع محدودی داشته باشد، فرصت ارتکاب تقلب را شناسایی می‌کند. ادراک فرصت‌ها ممکن است مدیریت بی‌پروا یا فقدان کنترل‌های داخلی در یک واحد تجاری باشد که تقلب را به یک فعالیت آسان تبدیل می‌کند. در نهایت، فرد تصمیم خود برای ارتکاب تقلب با متقاعد کردن خود مبنی بر اینکه به پول بیشتری نیاز دارد یا در نهایت آن را پس می‌دهد، منطقی می‌کند. با توجه به شرایط اقتصادی فشارها و انگیزه‌هایی برای ارتکاب کلاهبرداری افزایش یافته است که این امر باعث می‌شود کلاهبرداران به راحتی

اقدامات خود را منطقی کنند. تقلب در صنعت بانکداری رایج است و شامل؛ فیشینگ ایمیل، کلاهبرداری از کارت اعتباری، پول‌شویی، تقلب در درخواست وام، تقلب در صورت‌های مالی و کلاهبرداری سایبری می‌شود. با ظهور بانکداری دیجیتال، کلاهبرداری دیجیتال نیز در این بخش رایج‌تر شده است. بنابراین، مهم است که اذعان کنیم که مدیریت تقلب در صنعت بانکداری و بازگانی ضروری شده است، که مسلماً فرآیندی طاقت‌فرسا است. کلاهبرداران در کشف حفره‌ها ماهر شده‌اند و تکنیک‌های مؤثری مانند فیشینگ برای افراد ناآگاه و کلاهبرداری خلاقانه از آنها ایجاد کرده‌اند (ماریوت^۴، ۲۰۲۱). بنابراین، روش‌های کشف کلاهبرداری باید به طور مداوم تکامل یابند زیرا کلاهبرداران در طراحی تکنیک‌هایی مؤثرتر می‌شوند که سیستم‌های امنیتی سفت و سخت بانکی را دور می‌زنند و یاد می‌گیرند که چگونه افراد ناآگاه را متقاعد کنند که پول خود را در اختیار آنها بگذارند. تقلب مربوط به پرداخت یکی از جنبه‌های کلیدی آژانس‌های جرائم سایبری است و تحقیقات اخیر نشان داده است که استفاده از تکنیک‌های یادگیری ماشین برای شناسایی تراکنش‌های جعلی در مقادیر زیادی از داده‌های پرداخت مؤثر است. چنین تکنیک‌هایی توانایی شناسایی معاملات متقلبانه‌ای را دارند که حساب‌برسان انسانی ممکن است نتوانند آن‌ها را شناسایی کنند و همچنین این کار را به صورت بلادرنگ انجام دهند (نیکولاس، کوپا و لیخاک^۵، ۲۰۲۱). کشف دقیق تقلب^۶ (FD) می‌تواند اعتماد مشتریان بانک را در انجام امور بانکی‌شان به صورت آنلاین افزایش دهد. با پیشرفت تکنولوژی، تقلب به طور چشمگیری افزایش می‌یابد، که منجر به زیان‌های قابل توجهی برای شرکت‌ها می‌شود، بنابراین شناسایی تقلب به مسئله‌ای حائز اهمیت تبدیل شده است (کو، لو، سیوانگتان و هانگ^۷، ۲۰۲۰). بررسی شاخص ایمنی رایانه مایکروسافت نشان داد که تأثیر سالانه فیشینگ و اشکال مختلف سرقت هویت می‌تواند

1- Williams

2- Shihembetsa

3- CURTIS

4- Maruti

5- Nicholls, Kuppa, & Le-Khac

6- Fraud detection

7- Kou, Lu, Sirwongwattana, & Huang

غیرقابل پیش بینی ممکن است اشتباه کند. بنابراین، دقت کمتری دارد، زیرا همیشه به دانش فعلی خود اعتماد دارد؛ حتی برای مواردی که دانش کافی نیست.

کشف قلب برای بانکداری آنلاین یک زمینه مطالعاتی بسیار مهم است، زیرا مجرمان سایبری حملات جدید کلاهبرداری پیچیده‌ای را به صورت روزانه طراحی می‌کنند، بنابراین این امر مستلزم آن است که محققان تکنیک‌های جدید کشف قلب را به طور مداوم توسعه دهند. طبق نظر (ماراتونا^۴، ۲۰۱۳) در دسترس بودن اطلاعات دقیق در مورد سیستم کشف قلب بسیار محدود است زیرا صنعت بانکداری به ندرت آمار کشف قلب را منتشر می‌کند. به ویژه، تأمین کنندگان امنیت مؤسسات مالی، شرکت‌های شخص ثالثی هستند که از مالکیت معنوی خود در برابر رقبا محافظت می‌کنند. بنابراین، هم بانک‌ها و هم آژانس‌های فناوری اطلاعات؛ اطلاعات سیستم‌های امنیتی خود را منتشر نمی‌کنند (بولتون و هند^۵، ۲۰۰۲). همچنین، توسعه روش‌های جدید برای شناسایی قلب دشوار است زیرا تبادل افکار در مورد شناسایی قلب بسیار محدود است، اما نویسندگان بر این مفهوم تأکید می‌کنند که تکنیک‌های کشف قلب نباید به طور عمومی بیان شود. در غیر این صورت مجرمان ممکن است از همان داده‌ها بهره بگیرند (فوآ، لی، اسمیت و گایلر^۶، ۲۰۱۰). اغلب دو انتقاد اصلی از داده‌کاوی برای شناسایی قلب وجود دارد: کمبود اطلاعات تجربی واقعی در دسترس عموم، و فقدان تکنیک‌ها و روش‌های به خوبی تبلیغ شده است.

در دو دهه گذشته، با تکامل فناوری مورد استفاده در بخش بانکداری مالی، طرح‌های قلب مورد استفاده کلاهبرداران نیز تغییر کرد (استوجنا^۷، ۲۰۱۹). دو حوزه اصلی که در آن قلب صورت می‌گیرد، شامل؛ برنامه‌های کاربردی وب یا موبایل بانکداری اینترنتی و پرداخت‌های ATM، POS یا تاجر آنلاین با استفاده از کارت‌های بانکی می‌باشد. گزارش نیلسون^۸ (۲۰۲۰). بر افزایش هدف قرار دادن بازگازان

به ۵ میلیارد دلار آمریکا برسد، درحالی‌که هزینه ترمیم آسیب به شهرت افراد آنلاین می‌تواند تا ۶ میلیارد دلار یا بیشتر باشد. به طور متوسط برای هر ضرر ۶۳۲ دلار آمریکا تخمین زده می‌شود (مارکیان و لیم^۱، ۲۰۱۴). روش‌های سنتی تشخیص قلب در صنعت بانکداری مبتنی بر قانون بوده است که در آن انسان‌ها قوانین را تعریف می‌کنند. ۹۰ درصد مؤسسات مالی و بانکی بر این روش‌ها تکیه می‌کنند. درحالی‌که افراد بیشتری از فناوری‌های جدید استفاده می‌کنند، سناریوهای قلب بیشتری ممکن است؛ اتفاق بیفتد و این روش‌های مبتنی بر قوانین را در آینده غیرقابل مهر و موم و ناپایدار می‌سازد. علاوه بر این، مثبت کاذب (یعنی تراکنش‌های غیرمقلبانه که به عنوان جعلی فهرست‌بندی شده‌اند) باعث ضرر میلیون‌ها دلاری در معاملات و شکایات مشتریان در صنعت بانکداری می‌شود. روش‌های مبتنی بر قانون کمک زیادی به این نتایج می‌کنند. کیوبانو^۲ (۲۰۲۰) مطالعه‌ای را با ۱۰۰۰ مصرف‌کننده بزرگسال انجام داد و در آن متوجه شد حدود ۲۵ درصد از آنها که تراکنش‌هایشان رد شده بود، به تجارت با رقبا تمایل داشتند. این میزان روی آوردن به رقبا برای مصرف‌کنندگان ۱۸ تا ۲۴ ساله به ۳۶ درصد افزایش یافته است. همچنین برای افراد بین ۲۵ تا ۳۴ سال به ۳۱ درصد افزایش یافته است. این نتایج مطالعه نشان‌دهنده نیاز مبرم به روش‌های دقیق‌تر و مدرن‌تر کشف قلب است. به جای جلوگیری از تراکنش غیرمجاز، سیستم‌های بانکداری اینترنتی کشف قلب باید بلافاصله کلاهبرداری‌ها را در یک حساب در معرض خطر شناسایی کنند (ریچاردز^۳، ۲۰۰۳). رویکرد مبتنی بر قوانین متعارف برای کسب دانش (KA) برای یک تجارت بسیار کند، کار فشرده و پرهزینه‌ای است. شکنندگی نیز یکی از کاستی‌های پایه‌های قوانین مرسوم بود و این سیستم‌ها همیشه سعی می‌کنند پاسخی بدهند، حتی ممکن است نادرست باشد. شکنندگی به نرم‌افزاری اطلاق می‌شود که در مواجهه با یک موقعیت

- 1- Marican, & Lim
- 2- Ciobanu
- 3- Richards
- 4- Maruatona
- 5- Bolton, & Hand
- 6- Phua, Lee, Smith, & Gayler
- 7- Stojanov
- 8- Nilson Report

دارد. عملی‌ترین گزینه الگوریتم‌های یادگیری ماشین نصب شده در سیستم‌های بانکی هستند (گابینو و همکاران، ۲۰۲۱). با توجه به ضعف‌های ذاتی رویکرد تشخیص تقلب مبتنی بر قانون در بانک‌ها و داده‌های محدودی که بر الگوریتم‌های یادگیری ماشین تحت نظارت متداول استفاده می‌شود، نیاز مبرمی به تکنیک‌ها یا سیستم‌های جدید شناسایی وجود دارد که بتوانند با افزایش سریع تقلب‌ها و موارد پول‌شویی مقابله کنند. لذا این مقاله با هدف بررسی روش‌های کشف تقلب بانکی، به مرور ادبیات اخیر مرتبط با این موضوع می‌پردازد.

۲- مبانی نظری

ادبیات تخصصی و راه‌حل‌های موجود در بازار در دهه گذشته بر جمع‌آوری مقادیر قابل‌توجهی از داده‌ها در مورد تراکنش‌ها (و رفتار کاربر) و اصلاح الگوریتم‌های مورد استفاده برای شناسایی تقلب متمرکز شده‌اند. در همین راستا، قوانین اتحادیه اروپا (به‌عنوان مثال، PSD۲) به‌منظور تحمیل تعهدات به ذینفعان برای شناسایی تقلب به تصویب رسیده است. با این حال، از یک سو قانون تشریح سطح بالایی از این الزام قانونی ارائه می‌کند و از سوی دیگر، راه‌حل‌های موجود در بازار از نظر داده‌های جمع‌آوری‌شده و به‌ویژه تلاش برای تجمیع داده‌ها به‌منظور تولید متنوع می‌باشد. نتایج دقیق‌تر منجر به مبحثی می‌شود که هنوز در ادبیات تخصصی یا توسط قانون‌گذاران به‌طور عمیق مورد تجزیه و تحلیل قرار نگرفته است. محققانی مانند کارمیناتی و همکاران^۵ (۲۰۱۸)، در دهه گذشته شیوه‌های مختلفی را که می‌توان از چنین جزئیاتی برای کشف و پیشگیری از تقلب استفاده کرد، تحلیل کرده‌اند (کارمیناتی، و همکاران^۶، ۲۰۱۸). نتیجه‌گیری‌های آن‌ها منجر به رویکردهای مختلفی برای الگوریتم‌های کشف تقلب، با تأکید بر روش‌های یادگیری ماشین شده است (بائو، هیلاری، و کی^۷، ۲۰۲۲)؛ زیرا تجمیع داده‌ها و تحلیل

توسط سازمان‌های جرائم مالی سازمان‌یافته برای ارتکاب کلاهبرداری، با توسعه فناوری اطلاعات کشور و بازرگان که بر توانایی تاجر برای جلوگیری از تقلب تأثیر می‌گذارد، تأکید کرده است (گابینو، بریک، ماری، میهای و اسچو^۱، ۲۰۲۱؛ ناتان، ویکتور، گان و کات^۲، ۲۰۱۹). حدود ۵۶ درصد از اروپایی‌ها نگران این هستند که قربانی تقلب شوند (گابینو، و همکاران، ۲۰۲۱). در سال ۲۰۱۹، ۲۶ درصد از جمعیت اتحادیه اروپا گزارش دادند که پیام‌های جعلی دریافت کرده‌اند، از جمله پیام‌های مربوط به اعتبار بانکداری الکترونیک (اروستا و کامپوننت^۳، ۲۰۲۰). خانواده‌های مختلف بدافزار آسیب‌های مختلفی به مصرف‌کننده، زیرساخت‌های حیاتی، مؤسسات مالی و بانکی وارد کرده‌اند و به اهداف مورد نظر تبدیل شده‌اند (اسچو، گافتیا، آچیم، و کوتاک^۴، ۲۰۲۰). مکانیسم‌های تقلب از نظر بانکداری اینترنتی و تراکنش‌های کارت دارای ویژگی‌هایی هستند که می‌تواند به کشف تقلب کمک کند، مانند مکان شروع پرداخت، جزئیات گیرنده پرداخت، مهر زمانی پرداخت.

برای افزودن به چالش سیستم سنتی مبتنی بر قوانین، کلاهبرداران فاقد الگوهای خاص هستند و دائماً رفتار خود را در طول زمان تغییر می‌دهند و سیستم‌ها را دست‌وپاگیر و به سرعت منسوخ می‌کنند. نیاز آشکار به تغییر رویکرد در سیستم‌های امنیتی در سیستم‌های بانکی وجود دارد. طبق گزارش نیلسون (۲۰۲۰)، پیش‌بینی می‌شد که کلاهبرداری مربوط به کارت‌ها تا سال ۲۰۲۰ تنها به مبلغ حیرت‌انگیز ۳۰ میلیارد دلار در سطح جهان برسد. علاوه بر این، با اختلال فناوری در بخش بانکی به دلیل وجود کانال‌های پرداخت متعدد مانند کارت‌های اعتباری و نقدی، گوشی‌های هوشمند، نرخ تراکنش‌ها طی چند سال گذشته به‌طور تصاعدی افزایش یافته است. کلاهبرداران همچنین تاکتیک‌های کلاهبرداری بسیار مؤثری را توسعه داده‌اند. با توجه به این وضعیت، نیاز به توسعه رویکردهای سخت و قوی‌تر کشف تقلب در بانک‌ها وجود

1- Gbudeanu, Brici, Mare, Mihai, & cheau

2- Nathan, Victor, Gan, & Kot

3- Eurostat, & Components

4- cheau, Gaftea, Achim, & Cotoc

5- Carminati

6- Carminati, et al

7- Bao, Hilary, & Ke

توجه به انواع الگورتم‌های تشخیص تقلب پیشنهاد شده در مقالات تحقیقاتی منتشر شده در سه سال گذشته (۲۰۱۸-۲۰۲۱)، موجود در پنج پایگاه داده تحقیقاتی (ACM، Science Direct، Emerald Full text، IEEE Transactions Springer-Link)، بر اساس کلیدواژه‌های خاصی با هدف شناسایی جنبه‌های حریم خصوصی در نظر گرفته شده توسط این الگورتم‌ها ("تشخیص تقلب بانکی GDPR"، "تشخیص تقلب بانکی حریم خصوصی"، "تکنیک‌های تشخیص تقلب بانکی تجمعی"، "تکنیک‌های تشخیص تقلب بانکی")؛ انواع الگورتم‌های تشخیص تقلب بانکی در جدول شماره (۱) ارائه شده است.

تاریخی داده‌ها می‌تواند به یافتن الگوهای تقلب کمک کند (پولیتو، آلیپس، و پاتساکیتس^۱، ۲۰۱۹). این الگوهای تقلب می‌توانند به شناسایی یا پیش‌بینی تراکنش‌های تقلبی احتمالی کمک کنند. الگورتم‌های تشخیص، ویژگی‌های موجود تقلب‌ها را با تراکنش‌های فعلی در حال تجزیه و تحلیل مطابقت می‌دهند، در حالی که الگورتم‌های پیش‌بینی تلاش می‌کنند تا کلاه‌برداری‌هایی را شناسایی کنند که ویژگی‌های متفاوتی نسبت به تقلب‌های تاریخی دارند. محققان عموماً بر دقت نتایج و افزایش جمع‌آوری داده‌ها و تجمیع داده‌ها برای دستیابی به این هدف، هم برای الگورتم‌های کارآگاهی و هم برای الگورتم‌های پیش‌بینی تمرکز کرده‌اند (جها، گالین، و وستلند^۲، ۲۰۱۲). با

جدول ۱- تعداد مقالات منتشر شده در سه سال اخیر در خصوص انواع الگورتم‌های تشخیص تقلب بانکی

پایگاه داده / عبارت کلیدی	ژورنال‌های Springer-Link	IEEE	Emerald	Science Direct	کتابخانه دیجیتال ACM
تشخیص تقلب بانکی GDPR	۱۶۷	۷۲	۲۱	۱۱۳	۵۵۲
تشخیص تقلب در بانکداری حریم خصوصی	۸۹۱	۵۵۴	۱۹۳	۲۵۶	۵۹۱
تکنیک‌های تشخیص تقلب بانکی تجمعی	۲۸۸	۱۸۱	۴۴	۹۰	۷۷۶

1- Politou, Alepis, & Patsakis

2- Jha, Guillen, & Westland

الگوریتم‌های یادگیری ماشین، تشخیص تقلب در کارت اعتباری با استفاده از روش‌های Pipeling و Ensemble Learning (باگا، گویال، گوپتا و گویال^۵، ۲۰۲۰)، و تشخیص تقلب در کارت اعتباری با استفاده از شبکه عصبی مصنوعی (پینتو، و سابریرو^۶، ۲۰۲۲). ابزارهای مورد استفاده برای شناسایی فعالیت‌های متقلبان برای الگوهای تقلب شناسایی شده، سال به سال پیچیده‌تر و کارآمدتر می‌شوند و الگوریتم‌های یادگیری ماشین یکی از بهترین گزینه‌ها در این زمینه هستند، زیرا پیشرفت‌های فناوری از مرزها عبور می‌کنند و موارد بیشتری در دسترس قرار می‌گیرند (بیلگین، مارکو و دمیر^۷، ۲۰۱۲). اوگرک، ام، اوگرک، ای. و باهتیر^۸ (۲۰۱۹) سعی کردند روش‌های تقلب کارت اعتباری را با استفاده از مدلی با مجموعه داده Kaggle ارزیابی کنند. روش انتخاب شده شبکه‌های عصبی مصنوعی چند لایه^۹ (MANN) بود. برای شناسایی تقلب، این محققین از ویژگی‌هایی مانند: نقدی/خروجی، بدهی، پرداخت، مبلغ تراکنش یا ارزش محلی استفاده کرده‌اند. نتایج مطالعه نشان داد که روش انتخاب شده مؤثر است. همچنین، لی، هانگ، لیو و جیانگ^{۱۰} (۲۰۲۱)، با استفاده از مجموعه داده Kaggle، ایده یک روش ترکیبی را برای کاهش عدم تعادل طبقاتی با همپوشانی بر اساس ایده تقسیم و فتح ارائه کردند. برای این منظور از معیار ارزیابی آنتروپی وزن دار پویا استفاده شد. این آزمایش حتی از آزمایش‌های قبلی موفق‌تر بود مطالعاتی که رفتار افراد در انجام تراکنش‌ها را به عنوان روشی برای شناسایی تراکنش‌های جعلی کارت هدف قرار می‌دهند نیز وجود دارد: کارمیناتی و همکاران (۲۰۱۸) Banksealer را به عنوان یک سیستم پشتیبانی تصمیم ارائه کردند. این سیستم، برای هر کاربر یک مدل می‌سازد و سپس آن را بر روی این سیستم مدل می‌کند. در مرحله دوم، استحکام سیستم Banksealer در برابر مجرمان احتمالی بررسی شد. این رویکرد در ایتالیا پیاده‌سازی شد،

از سال ۲۰۱۸، مطالعات مربوط به سیاست حفظ حریم خصوصی داده‌های شخصی^۱ (سیاست GDPR) شروع به ظاهر شدن کردند. استالابوردیلون، پیارس و تسکالاسکی^۲ (۲۰۱۸) یک تحلیل بین‌رشته‌ای از GDPR در زمینه طرح‌های شناسایی الکترونیکی انجام داد. این مطالعه در بریتانیا انجام شد و یک نمای کلی از نحوه تفسیر این موقعیت‌ها ارائه کرد. در این مطالعه، نویسنده یک مبنای قانونی را پیشنهاد می‌کند که می‌تواند به مدیریت خوب حریم خصوصی توسط هر دو طرف درگیر کمک کند. علاوه بر این، در ارتباط با این موضوع، مطالعاتی در مورد میزان نفوذ در درخواست داده‌های شخصی شروع شده است. در این رابطه هوراک، استاپکا و هوساک^۳ (۲۰۱۹) به مسائل مربوط به تأثیر GDPR بر نرم‌افزار امنیت سایبری، از نظر عملیاتی پیشگیری از حادثه و رسیدگی به حوادث توجه کردند. خطرات نقض محرمانه بودن و اصل به حداقل رساندن داده‌ها با درخواست داده‌های شخصی مورد بررسی قرار گرفت، و همچنین این واقعیت که به اشتراک‌گذاری این اطلاعات توسط مشتری می‌تواند نگرانی‌های مشابهی را ایجاد کند. ارزیابی تأثیر حریم خصوصی داده‌ها^۴ (DPIA) که برای این سناریو انجام شد، نشان داد که با توجه به مکانیسم‌های کاهش خطر خاص، ریسک‌ها بالا نیستند. روش مورد استفاده در این مقاله به درک ریسک‌های موجود و مدیریت آسان‌تر آنها کمک کرد.

این مسائل مزاحمت حریم خصوصی نیز ارتباط نزدیکی با جلوگیری از تقلب کارت اعتباری دارد. تعداد قابل توجهی از نویسندگان به موضوع کشف تقلب پرداخته‌اند و الگوریتم‌های مورد استفاده به منظور ترکیب مکانیسم‌ها و طرح‌های جدید مورد استفاده توسط کلاه‌برداران به طور مداوم در حال بهبود هستند. تعدادی از مطالعات وجود دارد که تقلب کارت را از دیدگاه‌های مختلف تجزیه و تحلیل می‌کند: تشخیص تقلب کارت اعتباری با استفاده از

1- General data protection legislation

2- Stalla- Bourdillon, Pearce, & Tsakalakis

3- HorakHorák, Stupka, & Husák

4- Data Privacy Impact Assessment

5- Bagga, Goyal, Gupta, & Goyal

6- Pinto, & Sobreiro

7- Bilgin, MARCO LAU, & Demir

8- Örek, M., Örek, E., & Bahtiyar

9- Multi-layered artificial neural networks

10- Li, Huang, Liu, & Jiang

اما نتايج مطالعه نشان داد كه مي‌تواند مجموعه داده مفيدی در سراسر جهان ارائه كند (كارمينيات، و همكاران^۱، ۲۰۱۸). چن و همكاران^۲ (۲۰۱۹) نيز روشی را براي تشخيص تراكنش‌ها بر اساس رفتار افراد ايجاد كردند. مدلی كه در مقاله آنها استفاده شد: هايپر كره^۳ نام دارد. اين مطالعه به اين نتيجه رسيد كه اثر توصيف رفتار انسان با فراواني معاملات آنها مرتبط است (چن و همكاران، ۲۰۱۹). علاوه بر اين، وانگ و وانگ^۴ (۲۰۱۹) رفتار كاربر را از منظر فاصله زمانی بين درخواست ارزيايی كردند و دريافتند رفتارهای ربات مانند، در سيستم بانكداري آنلاين وجود دارد. نتايج مطالعه نشان داد كه پس از مقايسه دو الگوريتم، آنترپي Renyi در تمايز رفتار ربات از رفتار انسان نسبت به آنترپي شانون^۵ برتری دارد. محققين ديگر از مدل‌های مختلفي به منظور شناسايی تراكنش‌های تقلبي استفاده كردند. چن و همكاران (۲۰۲۰) يك مدل امتيازدهی ترکیبی را براي اين منظور توسعه داد. اين مدل مي‌تواند يك امتياز اعتباری دقيق به دست آورد و حتی ريسك اعتباری را کاهش دهد (چن، ياداك، خان، و ژو، ۲۰۲۰). مصر، تاكار، گوش و ساها^۶ (۲۰۲۰) يك مدل دومرحله‌ای برای شناسايی تراكنش‌های تقلبي پيشنهاده كردند. در مرحله اول از يك رمزگذار خودكار استفاده می‌شود كه ویژگی‌های تراكنش را به يك بردار مشخصه كوچك‌تر تبديل می‌كند. در مرحله دوم از بردار به عنوان ورودی طبقه‌بندی كننده استفاده می‌شود. آزمایش بر روی يك مجموعه داده مقايسه‌ای انجام و مشاهده شد كه مدل دومرحله‌ای كارآمدتر از سيستم‌هایی است كه تنها با یکی از دو مرحله طراحی شده‌اند.

اولووكره و ادوال^۸ (۲۰۲۰) چارچوبی را پيشنهاده می‌كنند كه تكنيك‌های گروه فرا يادگیری و يادگیری حساس

به هزينه را براي كشف تقلب تركيب می‌كند. نتايج اين مطالعه نشان داد كه چارچوب مجموعه حساس به هزينه، طبقه‌بندی‌كننده‌های حساس به هزينه را توليد می‌كند كه در تشخيص تقلب در پايجاه‌های داده پرداخت‌ها كارآمد هستند. مطالعات ديگر بر روی تشخيص‌های خاص‌تر متمرکز شده‌اند. آماراسينگه، آپونسو، و كريشنارگيا^۹ (۲۰۱۸) يادگیری ماشين انتخاب شده و تكنيك‌های تشخيص قبلی را كه می‌توانند در يك سيستم تشخيص تراكنش مالی تقلبی ادغام شوند، بررسی كردند. نتيجه‌گیری شد كه به منظور شناسايی مؤثرتر تقلب بانکی، دانستن الگوريتم‌های خاص (به عنوان مثال، شبكه‌های بیزی، منطق فازی و غيره) مهم است. دونگ و همكاران^{۱۰} (۲۰۱۸) روی حوزه بسيار جالبی از موضوع تقلب، يعنی تقلب‌های تبليغاتی موبایلی كار كردند و يك رويكرد ترکیبی به نام FraudDroid را براي شناسايی تقلب در برنامه‌های کاربردی در دستگاه‌های اندرویدی پيشنهاده كردند. پس از تجزيه و تحليل ۱۲۰۰۰ برنامه مشكوك، FraudDroid ۳۳۵ مورد تقلب را شناسايی كرد و اين نتيجه تأييد می‌كند كه روش مفيدی برای كشف اين نوع جرم است. علاوه بر اين، برخی محققين، تحقق رأی از طريق دستگاه ATM با ارائه احراز هويت بيومتريك يا احراز هويت تشخيص چهره را پيشنهاده كردند. استفاده از برنامه رأی‌گیری ATM در مقايسه با كارت‌های Aadhar برای امنيت و حفظ حریم خصوصی، آسان‌تر است (سادهارسان، كومار، و تكاسان، ساتياپريا و سارانایا^{۱۱}، ۲۰۱۹)

چالش‌های بانكداري آنلاين

هر ساله تقلب در اشكال مختلف افزايش می‌يابد كه منجر به خسارات مالی قابل توجهی می‌شود (وی، لی، كائو،

- 1- Carminati, et al
- 2- Chen
- 3- Hypersphere
- 4- Wang and Wang
- 5- Shannon entropy
- 6- Chen, Yadav, Khan, & Zhu
- 7- Misra
- 8- Olowookere, & Adewale
- 9- Amarasinghe
- 10- Dong
- 11- Sudharsan, Kumar, Venkatesan, Sathyapreiya, & Saranya

بسیار محدود است و بانک‌ها اغلب آمار سیستم‌های FD را منتشر نمی‌کنند (اسکانسی، گانو^۷، ۲۰۱۷؛ ماراتون، ۲۰۱۳). زیرا بیشتر امنیت توسط شرکت‌های IT شخص ثالث ارائه می‌شود که از مالکیت معنوی در برابر رقبای خود نیز محافظت می‌کنند. بنابراین هم بانک‌ها و هم شرکت‌های امنیتی فناوری اطلاعات بیشتر اطلاعات سیستم‌های امنیتی خود را منتشر نمی‌کنند. بولتون و هند^۸ (۲۰۰۲) نیز تأکید می‌کنند که توسعه روش‌های جدید کشف تقلب دشوار است؛ زیرا تبادل نظر در کشف تقلب بسیار محدود است، اما از این ایده حمایت می‌کنند که تکنیک‌های FD نباید به‌طور عمومی با جزئیات توصیف شوند. در غیر این صورت مجرمان نیز ممکن است به آن اطلاعات دسترسی پیدا کنند (کارمیناتی و همکاران، ۲۰۱۵). فوا و همکاران^۹ (۲۰۱۰) تأکید می‌کنند که کشف تقلب با استفاده از تکنیک‌های داده‌کاوی در صنعت بسیار رایج است.

تقلب در سرمایه‌گذاری سهام و اوراق بهادار

بازار سهام و اوراق بهادار مالی به مردم این امکان را می‌دهد که پول خود را با هدف کسب بازدهی مثبت بر اساس انجام تحقیقات یا فقط یک حدس، سرمایه‌گذاری کنند. با این حال، مشخص است که بخشی از فعالان بازار تقلب می‌کنند و با این کار سودهای کلانی به قیمت سرمایه‌گذاران نهادی و خرد به دست می‌آورند. دستگیری این بازیگران کلاه‌بردار آسان نیست و معمولاً به نیروی کار زیادی برای جمع‌آوری شواهد در مدت زمان طولانی نیاز دارد، به‌ویژه در موارد تجارت داخلی. با این حال، پیشرفت‌های اخیر در کاربردها و تکنیک‌های یادگیری ماشین به شناسایی بازیگران بد به روشی کارآمدتر و سریع‌تر کمک می‌کند. برخی از روش‌های مورد استفاده توسط افرادی که فعالیت سرمایه‌گذاری متقلبانه انجام می‌دهند عبارتند از: دستکاری بازار، تجارت داخلی، پول‌شویی، تأمین مالی تروریسم و بسیاری موارد دیگر. دستکاری بازار

اووو چن^{۱۰} (۲۰۱۳). بر اساس بررسی شاخص ایمنی محاسباتی مایکروسافت (MCSI) در سال ۲۰۱۴، تأثیر سالانه فیشینگ و اشکال مختلف سرقت هویت در سراسر جهان می‌تواند تا ۵ میلیارد دلار آمریکا تخمین زده شود، در حالی که هزینه ترمیم آسیب به شهرت آنلاین افراد می‌تواند بیشتر از ۶ میلیارد دلار آمریکا، یا میانگین تخمینی ۶۳۲ دلار آمریکا به ازای هر ضرر باشد (ماریکان و لیم، ۲۰۱۴). مشاوره PwC در نظرسنجی جهانی جرم اقتصادی در سال ۲۰۱۶ اشاره می‌کند که تقریباً یک سوم شرکت‌های مورد بررسی گزارش داده‌اند که قربانی نوعی از جرائم سایبری بوده‌اند. همچنین، جرائم سایبری دومین نوع رایج جرائم اقتصادی است که در این نظرسنجی تجزیه و تحلیل شده است، در حالی که یک نظرسنجی اخیر گزارش می‌دهد (لاویون^{۱۱}، ۲۰۱۸). از سال ۲۰۱۶ بالغ بر ۱۳ درصد افزایش در تقلب رخ داده است. ۱۲۳ یک منبع ارزشمند برای قربانیان جرائم اینترنتی و مجریان قانون در شناسایی، رسیدگی و تعقیب جرائم است (اف.بی.آی، ۲۰۱۸). ۱۲۳ گزارش می‌دهد که ۱۴۴۰۸ شکایت در سال ۲۰۱۸ دریافت کرده که مربوط به کلاهبرداری پشتیبانی فنی از قربانیان از ۴۸ کشور بود. زیان گزارش شده نشان‌دهنده افزایش ۱۶۱ درصدی زیان نسبت به سال قبل می‌باشد. گزارش نرخ جرائم اقتصادی جهانی توسط (اسکانسی، گانو^{۱۲}، ۲۰۱۷)، نشان می‌دهد که خدمات مالی با جرم اقتصادی ۴۸ درصد بیشترین بخش در معرض خطر بوده و دومین جرم اقتصادی گزارش شده در جرائم سایبری است. رایج‌ترین دامنه‌های تقلب عبارتند از: بانکداری آنلاین، کارت‌های اعتباری، مخابرات، بیمه مراقبت‌های بهداشتی، بیمه آنلاین، نفوذ کامپیوتر، حراجی آنلاین (اسکانسی، گانو، ۲۰۱۷؛ بولتون و هند، ۲۰۰۲؛ جان، آئل، کندی، اولاجید و کندی^{۱۳}، ۲۰۱۶؛ عبدالله، معروف و زینال^{۱۴}، ۲۰۱۶؛ وی و همکاران، ۲۰۱۳). بنابراین، انجام تحقیقات کشف تقلب برای بانکداری آنلاین یک کار بسیار مهم است، اما چالش‌هایی در این زمینه وجود دارد که باید برطرف شوند. دانش مکانیزم کشف تقلب بانک‌ها

1- Wei, Li, Cao, Ou, & Chen

2- Lavion

3- Ocansey, & Ganu

4- John, Anele, Kennedy, Olajide, & Kennedy

5- Abdallah, Maarof, & Zainal

6- Ocansey, & Ganu

7- Bolton & Hand

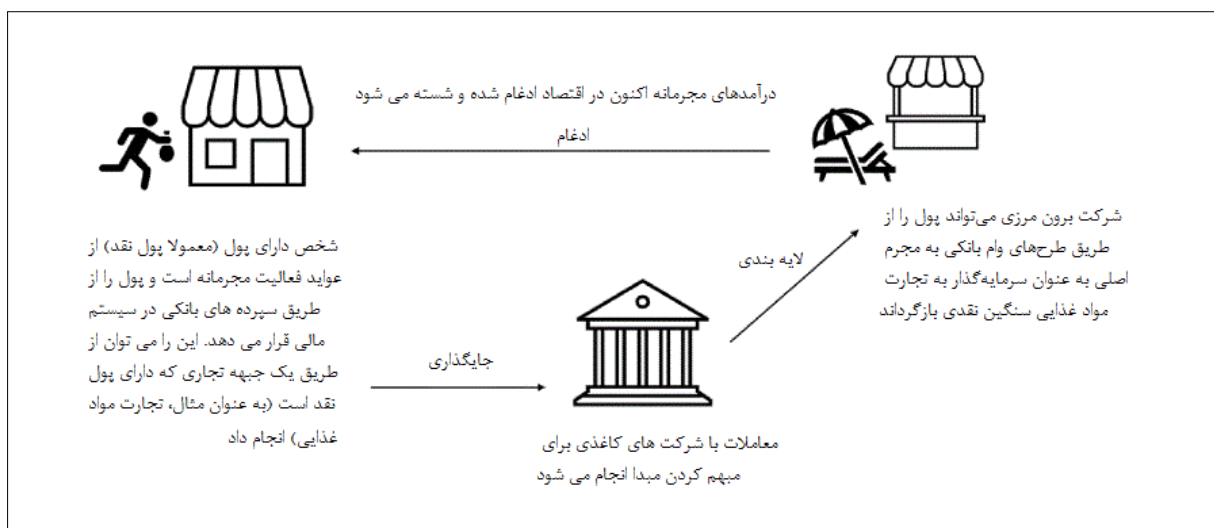
8- Phua

آوردن پول از راه دور استفاده شود. پول شویی روشی است که توسط مجرمان و افرادی که دارایی‌های «کثیف» یا به طور غیرقانونی از طریق فعالیت‌های مجرمانه به دست آمده‌اند، برای تبدیل این پول‌ها به یک حالت «پاک» یا مشروع از نظر قانون و دولت‌ها استفاده می‌کنند. سرویس دادستانی سلطنتی^۱ (CPS) بریتانیا طرح پول‌شویی را معمولاً شامل سه مرحله تعریف می‌کند: اول، جایگذاری است که فرآیند واریز پول مجرمانه به سیستم مالی است. دوم، لایه‌بندی است که پول را در سیستم مالی از طریق شبکه‌های پیچیده تراکنش‌ها با هدف مبهم کردن حرکت می‌دهد. لایه‌بندی معمولاً از طریق شرکت‌های خارج از کشور انجام می‌شود. درنهایت، ادغام، پول مجرمانه‌ای است که از طریق سرمایه‌گذاری‌هایی مانند املاک، خرید سهام و اقلام لوکس در اقتصاد واقعی جذب یا ترکیب می‌شود (هاپتون^۲، ۲۰۱۶). نمونه‌ای از پول‌شویی با استفاده از قرار دادن، لایه‌بندی و ادغام را می‌توان در شکل شماره (۱) مشاهده کرد.

اقدامی برای فروش یا خرید یک اوراق بهادار مالی با هدف دست‌کاری هدفمند قیمت دارایی یا اوراق بهادار پایه در نظر گرفته می‌شود. تجارت داخلی غیرقانونی زمانی است که «خودی‌ها» یا افرادی که از مطالب شرکت خصوصی و غیردولتی مطلع هستند، از آن اطلاعات قبل از انتشار عمومی برای بهره‌مندی مالی استفاده می‌کنند. این نه تنها شامل معامله اوراق بهادار، بلکه نشت اطلاعات غیر عمومی به اشخاص ثالث نیز می‌شود.

کشف تقلب و مبارزه با پول‌شویی (AML)

اغلب اوقات، بین تقلب بانکی و سرقت از بانک سردرگمی وجود دارد. دزدی از بانک با تقلب بانکی متفاوت است، چرا که سرقت از بانک شامل خشونت می‌شود، درحالی‌که تقلب بانکی اغلب یک فرآیند آهسته و مخفیانه است که تا زمانی که عمل کامل شود، شناسایی نمی‌شود. تقلب بانکی اغلب مستلزم ترغیب مشتریان به ارائه اطلاعات حساس بانکی خود است که ممکن است برای بدست



شکل (۱): مراحل پول‌شویی

1- Crown Prosecution Service
2- Hopton

كلاهبرداری در بانكداري اينترنتي، بانكداري تلفن همراه، فيشينگ، كلاهبرداری در سايت خريد و حراجي، كلاهبرداری، هرزنامه، و سرقت هويت انواع مختلفی از كلاهبرداری‌های آنلاين هستند (ای.پی.اف. ۲۰۱۸). تقلب بانكداري اينترنتي نوعی كلاهبرداری است كه با استفاده از هر فناوری آنلاين برای برداشت غيرقانونی پول یا انتقال آن به حساب بانكي دیگر انجام می‌شود. جرم مالی یا جرائم اقتصادی توسط يوروپل به عنوان «اعمال غيرقانونی انجام شده توسط يك فرد یا گروهی از افراد برای به دست آوردن مزیت مالی یا حرفه‌ای» تعریف شده است (نيكولس، و همكاران، ۲۰۲۱). انگیزه اصلی در چنین جنایاتی، منفعت اقتصادی است. مجمع جهانی اقتصاد، جرائم مالی را صنعتی تريليون دلاری می‌داند (ناندوری، جیا، اوکا، بیوار و لیو، ۲۰۲۰). جرائم مالی در فضای مجازی با استفاده از ابزارهای هك و تكنیک‌های مهندسی اجتماعی كه امنیت مؤسسات مالی و شركتی را دور می‌زند، انجام می‌شود. این امر باعث می‌شود كه محققان و شركت‌ها به جرائم مالی از منظر دیگری نگاه كنند، به این ترتیب كه تمایز بین جرائم مالی، هك و مهندسی اجتماعی برای منافع اقتصادی كم‌رنگ شده است. اینجاست كه نویسندگان اصطلاح جرائم سایبری مالی را معرفی می‌كنند كه تركیبي از جرائم مالی، هك و مهندسی اجتماعی است كه در فضای سایبری تنها با هدف كسب سود اقتصادی غيرقانونی انجام می‌شود. كشف تقلب در سازمان‌های مالی سنتی^۱ بیشتر طول می‌كشد (پاسكودال، مارچینی، و میلر^۲، ۲۰۱۷). این مقدار زمان برای مشتریان و مؤسسات مالی نامطلوب است؛ زیرا ممكن است چندین فعالیت متقلبانه قبل از شناسایی در چنین چارچوب زمانی رخ دهد. كلاهبرداری بر بانك‌هایی كه با خدمات پرداخت آنلاين سروكار دارند، به‌ویژه در پیشرفت‌های تکنولوژیک معاصر در بخش تجاری، تأثیر منفی می‌گذارد. به عنوان مثال، حدود ۲۰ درصد از مشتریان پس از تجربه كلاهبرداری، بانك را تغییر می‌دهند (ساندو^۳، ۲۰۲۱). این تعداد مشتریانی كه از يك بانك خاص خارج می‌شوند، به طور قابل توجهی برای عملیات تجاری يك

كشف تقلب مستلزم فعالیت‌هایی است كه از به دست آوردن پول یا دارایی از طریق جعل نادرست جلوگیری می‌شود. در بخش بانكي، تقلب مستلزم جعل چك و استفاده از كارت اعتباری سرقت شده است. ممكن است متضمن اغراق زیان یا ایجاد رویدادهای ناگوار مانند تصادفات با هدف صرف پرداخت باشد (شیهم بتسا، ۲۰۲۱). بر اساس گزارش نیلسون (۲۰۲۰)، بیش بینی می‌شد كه كلاهبرداری مربوط به كارت‌ها تا سال ۲۰۲۰ تنها به مبلغ حیرت‌انگیز ۳۰ میلیارد دلار در سطح جهان برسد. علاوه بر این، اختلال در فناوری در بانكداري و پرداخت‌ها به دلیل افزایش كانال‌های پرداخت مانند اعتبار و كارت‌های نقدی و گوشی‌های هوشمند، تعداد تراكنش‌ها طی سال‌های اخیر افزایش چشمگیری داشته است (گابینو و همكاران، ۲۰۲۱). تلاش‌های كشف تقلب با استفاده از تجزیه و تحلیل داده‌ها، نرم‌افزار طراحی‌شده برای كشف تقلب، و ابزارها، برنامه‌های طراحی‌شده برای كشف تقلب، سازمان‌ها را قادر می‌سازد تا تكنیک‌های رایج كلاهبرداری را پیش بینی كنند، فرآیند ارجاع متقابل داده‌ها را خودكار كنند، تراكنش‌های خود را به‌طور مستمر در زمان واقعی نظارت كنند، و كشف كنند (ویلیام، ۲۰۲۲). منابع تشخیص و پیشگیری از تقلب مانند نرم‌افزار در دو نسخه اختصاصی و رایگان موجود است. برخی از ویژگی‌های رایج در این نرم‌افزار عبارتند از: داشبورد، واردات و صادرات داده‌ها، تجسم داده‌ها، یکپارچه‌سازی مدیریت ارتباط با مشتری، مدیریت تقویم، زمان‌بندی، بودجه‌بندی و مدیریت رمز عبور. آنها همچنین دارای رابط‌های برنامه‌نویسی کاربردی (API)، صدور صورت حساب، احراز هويت دومرحله‌ای و مدیریت پایگاه داده مشتری هستند.

سیستم‌های تشخیص تقلب در بانكداري

كلاهبرداری یا تقلب آنلاين، تقلب اينترنتي و جرائم سایبری مبنای گسترده‌ای دارند و به روش‌های مختلفی رخ می‌دهند. كلاهبرداری آنلاين را می‌توان به عنوان هر عمل غيرقانونی انجام شده به صورت آنلاين توصیف كرد.

- 1- Nanduri, Jia, Oka, Beaver, & Liu
- 2- brick-and-mortar financial organizations
- 3- Pascual, Marchini, & Miller
- 4- Sando

پول شويي توسط PRM شناسايي مي شود (حفيظ، عقيلي، و زاواراسكي، ۲۰۱۶).

سيستم SAS Fraud Manager: اين سيستم عمدتاً راه حل كشف قلب براي كارت هاي بدهي و اعتباري است و قابليت بلادرنگ را دارد. اين سيستم كشف قلب از تكنيك تشخيص ناهنجاري^۵ (AD) استفاده مي كند (اس. آي. اس. ۲۰۰۷).

سيستم Kount: يكي ديگر از سيستم هاي پيشگيري از قلب، كانت است كه از هر دو مدل ML^۶ نظارت شده و بدون نظارتي استفاده مي كند. برخي از بزرگ ترين ارائه دهندگان خدمات پرداخت، دروازه ها، كيف پول ها و پردازنده ها از اين سيستم استفاده مي كنند. تعداد كمی از سيستم هاي كشف قلب و پيشگيري ديگر وجود دارند (پاتم سي، ۲۰۱۵؛ فاندنت، ۱۹۹۷؛ ريسك نت، ۱۹۹۸).

- نقاط ضعف سيستم هاي تشخيص قلب تجاري

سيستم هاي كشف قلب تجاري ذكر شده در بالا به طور گسترده توسط بانك ها و مؤسسات مالي استفاده مي شوند. با اين حال، هيچ سيستم كشف قلب به طور ۱۰۰ درصد مؤثر نيست. برخي از كاستي ها در سيستم هاي كشف قلب تجاري عبارتند از:

اين سيستم ها فاقد امتيازدهي و رد يابي مكان مصرف كننده براي تراكنش دستگاه تلفن همراه هستند. لايه يکپارچه سازي براي اين سيستم ها براي وارد كردن داده ها كامل نيست.

همه سيستم ها قابليت ثبت رمزگذاري داده ها را ندارند.

سيستم ها مبالغ بيشتر از تراكنش ها را در نظر مي گيرند و عمدتاً مبالغ كمتر را ناديدده مي گيرند (ديون^۷، ۲۰۱۴؛ هرسچل، ليدن، و كارت^۸، ۲۰۱۵).

بانك مضر است، به خصوص اگر اين روند طی چندین سال ادامه یابد. ضرورت ایجاد رویکردهای مناسب و قوی كشف قلب در سيستم هاي مؤسسات مالي براي مهار اين عمل وجود دارد. دو رويکرد اصلي كشف قلب وجود دارد كه شامل رويکرد مبتني بر قانون و تشخيص قلب يادگيري ماشين است. سيستم هاي بانكداري آنلاين^۱ (OBS) مكانيسم هاي امنيتي مختلفي براي جلوگيري از دسترسي غيرمجاز كلاه برداران دارند. عليرغم همه اين پيشرفت ها، قربانيان بي خبر بازهم اعتبار خود را در برابر كلاه برداران فيشينگ و سارقان هويت آنلاين از دست مي دهند. هنگامي كه يك كلاه بردار به حساب بانكي کاربر دسترسي پيدا مي كند، تشخيص آن فعاليت آسان نيست. بانك هاي مختلف از سيستم هاي كشف قلب مختلفي استفاده مي كنند. برخي از سيستم هاي كشف قلب شناخته شده تجاري مورد استفاده براي بانكداري آنلاين: SAS Fraud Manager و FICO، PRM، Kount هستند.

سيستم FICO: سيستم FICO يكي از پرکاربردترین سيستم هاي كشف قلب مورد استفاده توسط بانك ها در سطح جهان است (اف. آي. سي. او.، ۲۰۱۰؛ باربازون، كاهيل و والش^۲، ۲۰۱۰). اين سيستم، از يك شبكه عصبي و يك موتور مبتني بر قانون استفاده مي كند. قابليت هاي بلادرنگ دارد. FICO توسط مؤسسات مالي، بانك ها، اتحاديه هاي توليدي و اعتباري، عمدتاً براي كشف قلب بدهي، اعتبار، سپرده و پرداخت الكترونيكي استفاده مي شود (كاپتر^۳، ۲۰۱۹).

سيستم PRM : يكي ديگر از سيستم هاي كشف قلب پر کاربرد PRM است (آي. سي. آي.، ۲۰۱۱). PRM مانند FICO، از يك شبكه عصبي و يك رويکرد مبتني بر قوانين نيز استفاده مي كند. سيستم PRM در بيش از ۴۰ کشور از جمله هشت بانك از ۲۰ بانك برتر جهان استفاده مي شود. انواع قلب از كارت بدهي و اعتباري، قلب حساب و قلب

1- Online Banking Systems

2- Brabazon, Cahill, Keenan, & Walsh

3- Capterra

4- Hafiz, Aghili, & Zavarsky,

5- Anomaly detection

6- Machine learning

7- Dehaven

8- Herschel, Linden, & Kart

سیستم‌های مبتنی بر قانون (RBS)

سیستم‌های مبتنی بر قانون^۱ بخشی از گروه بزرگی از رویکردها هستند که سعی در مدل‌سازی تخصص انسان به نام سیستم‌های مبتنی بر دانش^۲ (KBS) دارند. KBS سنتی از جمله RBS دارای مسائل مشترک KA، شکنندگی^۳ و یادگیری افزایشی^۴ است. اصطلاح شکنندگی توسط لیگ (۲۰۰۷) ابداع شده است و به نرم‌افزاری اطلاق می‌شود که احتمالاً در مواجهه با سناریوی غیرمنتظره به نتیجه نادرستی می‌رسد. درحالی‌که یادگیری افزایشی یکی از راه‌حل‌های ممکن برای مشکل مقیاس‌پذیری است، که در آن داده‌ها در بخش‌هایی پردازش می‌شوند و سپس نتایج را برای کاهش استفاده از حافظه ترکیب می‌کنند (سید، لیو، و سانگ^۵، ۱۹۹۹). فعالیت‌های مربوط به قلب در حوزه مالی با کاوش سیگنال‌های سطحی و واضح قابل شناسایی هستند. با وجود غیرعادی بودن، تراکنش‌های بزرگ و آنهایی که در مکان‌های غیرمعمول اتفاق می‌افتند، باید تأیید قوی‌تری داشته باشند. سیستم‌های مبتنی بر قانون از الگوریتم‌هایی استفاده می‌کنند که موقعیت‌های مختلف کشف قلب را ارزیابی می‌کنند که به صورت دستی توسط تحلیلگران قلب نوشته شده‌اند. در حال حاضر، سیستم‌های حقوقی تقریباً از ۳۰٪ قانون مختلف برای تأیید معاملات استفاده می‌کنند. این توضیح می‌دهد که چرا استفاده از سیستم‌های مبتنی بر قانون یک فرآیند ساده است. این سیستم‌ها نیاز به تنظیم دستی سناریوها و موقعیت‌هایی دارند که قادر به تشخیص همبستگی‌های ضمنی نیستند. علاوه بر این، سیستم‌های مبتنی بر قانون اغلب از نرم‌افزار قدیمی استفاده می‌کنند که به سختی جریان داده‌های بلندرنج را از طریق سیستم‌هایی که در حوزه دیجیتال مهم هستند، پردازش می‌کنند. سیستم‌های مبتنی بر قوانین؛ به قوانین از پیش تعیین شده بستگی دارند تا تغییرات در رفتار را تشخیص دهند. این سیستم‌ها سفت و سخت هستند و قادر به انطباق

با صنایعی مانند خدمات مالی نیستند که برای غلبه بر چالش‌های مرتبط با شناسایی قلب نیاز به پلت‌فرم‌های سبک‌تر و چابک‌تری دارند (مون و کیم^۶، ۲۰۱۷). رویکردهای مبتنی بر قانون زمان بر هستند زیرا به تحلیل گران قلب نیاز دارند تا قوانین مورد استفاده برای کشف قلب را ایجاد کنند. این رویکردها همچنین به کاردستی و چندین فرآیند تأیید نیاز دارند که مانع از تجربه کاربر می‌شود. رویکردهای مبتنی بر قوانین، الگوهای قلب آشکار را شناسایی می‌کنند و بنابراین با تغییرات کلاه‌برداری که با تکامل کلاه‌برداران رخ می‌دهد، سازگار نیستند. از سوی دیگر، مدل‌های مبتنی بر قانون برای نگهداری مجموعه داده‌ها پرهزینه‌تر می‌شوند یا اندازه پایگاه مشتری گسترش می‌یابد.

تشخیص قلب مبتنی بر یادگیری ماشین (ML)

رویدادهای پنهانی در رفتار کاربر وجود دارد که ممکن است کاملاً مشهود نباشد و تراکنش‌های جعلی احتمالی را به نمایش بگذارد. یادگیری ماشین اجازه ایجاد الگوریتم‌هایی را می‌دهد که می‌توانند مجموعه داده‌های بزرگ‌تری را با چندین متغیر پردازش کنند و به شناسایی این همبستگی‌های پنهان بین رفتار اپراتور و احتمال فعالیت‌های متقلبانه کمک می‌کنند. سیستم‌های یادگیری ماشین بهتر از سیستم‌های مبتنی بر قوانین هستند، زیرا در پردازش داده‌ها سریع‌تر هستند و عملیات دستی کمتری دارند. به عنوان مثال، الگوریتم‌های هوشمند با تجزیه و تحلیل رفتار در کاهش مراحل تأیید مورد نیاز مطابقت دارند. شرکت‌هایی که با مقررات خدمات مالی سرو کار دارند در نظارت بر فعالیت‌های قلبی احتمالی درگیر هستند؛ آنها باید در مورد فعالیت‌های علامت‌گذاری شده با یکدیگر ارتباط برقرار کنند. ویلالوبوس^۷ و همکاران (۲۰۱۹) نمونه‌ای را توصیف می‌کنند که در آن یک نمونه اولیه یادگیری ماشین بر روی مجموعه داده‌ای برنامه‌ریزی شد

1- Rule-based systems

2- Knowledge-Based Systems

3- Brittleness

4- Incremental learning

5- Syed, Liu, & Sung

6- Moon, & Kim

7- Villalobos

میزان ایمنی بیشتری را ارائه دهند، اما نسبت به شدت ایمنی، به منابع محاسباتی بسیار بیشتری نیاز دارند (لی، پارک، اوم و چانگ^۲، ۲۰۱۱). نویسندگان یک IDS چند سطحی و مدیریت گزارش را برای IDS مؤثر در رایانش ابری پیشنهاد کردند. این امر با معرفی مقدار قابل توجهی از مسئولیت ایمنی به مشتریان بسته به نوع روش ناهنجاری، که کاربران را بر اساس نرخ ناهنجاری به سازمان‌های ایمنی متمایز متصل می‌کند، به استفاده کارآمد از منابع کمک می‌کند. برخی محققین، بر این باورند که هکرها می‌توانند یک سری از حملات را به یک سیستم مقصد امن در فضای ابری انجام دهند، به عنوان مثال، با فرار از یک دستگاه با استفاده آسان مبتنی بر ابر و سپس از درب پشتی^۳ قبلی برای حمله به سیستم استفاده می‌کنند (چن، گان، هانگ، و کیو^۴، ۲۰۱۲). سیستم تشخیص پیشنهادی لاگ‌های مختلف را از ابر برای به دست آوردن معانی فعالیت‌های گزارش تجزیه و تحلیل می‌کند. برای تعداد کمی از تخلفات، فعالیت‌های مشکوک اغلب توسط مدیر نادیده گرفته می‌شود. مدل پنهان مارکوف^۵ برای مدل‌سازی توالی حملات انجام‌شده توسط هکرها و چنین رخدادهای مخفیانه در یک چارچوب طولانی مدت اجرا می‌شود، زیرا این مدل آگاه از وضعیت می‌باشد. سیستم‌های پیشنهاد شده برای IDS، عمدتاً بر شناسایی رویدادهای بالقوه، ثبت داده‌ها و تلاش‌های نظارتی تمرکز دارند.

تشخیص ناهنجاری (AD)

تشخیص ناهنجاری (AD) شامل استفاده از تکنیک‌های مختلف محاسباتی و ریاضی برای تشخیص نقاط غیرعادی در یک مجموعه داده است. در ادبیات مرتبط، تشخیص ناهنجاری نام‌های مختلفی دارد: تشخیص بیرونی^۶، تشخیص تازگی^۷، تشخیص نویز^۸ و تشخیص انحراف^۹. تشخیص ناهنجاری فرآیند تجزیه و تحلیل مجموعه داده

که تراکنش‌های آن به صورت مجرمانه انجام شده بود. نمونه اولیه مورد استفاده در سیستم مبتنی بر قانون به کشف روابط پنهان بین معاملات و فعالیت‌های مجرمانه کمک کرد. چنین سیستم‌هایی حجم کار را در بانک‌های کوچک‌تر درگیر در نظارت بر تقلب به حداقل می‌رساند. راه حل پیشنهادی نشان داد که ۹۹٫۶ درصد از معاملات پول‌شویی و تراکنش‌های گزارش شده از ۳۰ درصد به ۱ درصد کاهش یافته است. ML به الگوریتم‌هایی بستگی دارد که با افزایش اندازه مجموعه داده‌ها مؤثرتر هستند. هر چه داده‌ها بیشتر باشد، نمونه اولیه ML بیشتر بهبود می‌یابد و می‌تواند شباهت‌ها و تفاوت‌ها را در رفتارهای مختلف تشخیص دهد. هرچه مدل ML تفاوت بین تراکنش‌های مشروع و جعلی را بیشتر شناسایی کند، سیستم‌ها در دسته‌بندی مجموعه‌های داده در دسته‌های مختلف کارآمدتر می‌شوند. بنابراین، سیستم‌های ML با رشد پایگاه داده مشتری، مقیاس‌پذیرتر هستند. در حالی که الگوریتم‌های ML مزایای متعددی را ارائه می‌کنند، اما دارای معایب قابل توجهی هستند که استفاده از آنها را در تشخیص تقلب محدود می‌کند. به عنوان مثال، یکی از اشکالات این است که ML برای دستیابی به دقت، به مقدار قابل توجهی داده نیاز دارد. این حجم داده قابل مدیریت است. با این حال، باید نقاط داده کافی وجود داشته باشد که روابط علی مشروع را در سازمان‌های کوچک‌تر تشخیص دهد. علاوه بر این، مدل‌های یادگیری ماشین بر روی اعمال، رفتار و فعالیت‌ها عمل می‌کنند. این مدل تمایل دارد اتصالات واضح را نادیده بگیرد، چنین کارتی در دو حساب مختلف استفاده می‌شود، از این رو، فعالیت کشف تقلب را بی‌اثر می‌کند.

-سیستم‌های تشخیص نفوذ (IDS)

با توجه به سیستم‌های تشخیص نفوذ^۱ (IDS) می‌توانند

1- Intrusion Detection Systems
2- Lee, Park, Eom, & Chung
3- Backdoor
4- Chen, Guan, Huang, & Ou
5- Hidden Markov Model
6- Outlier Detection
7- Novelty
8- Noise
9- Deviation

تکنیک AD اخطار را برای کاربران سیستم ارسال می‌کند و بیشتر متمرکز بر ابر است و برای اطلاعات بزرگ مناسب نیست (چیو، یه و لی^۱، ۲۰۱۳). برابازون و همکاران^۲ (۲۰۱۰) همچنین استفاده از یک رویکرد مبتنی بر AIS برای AD را برای کاهش تقلب در کارت اعتباری توصیف می‌کنند. تجزیه و تحلیل غیرعادی برای داده‌های گزارش جریانی توسط محققین انجام شده است (مروستی، پوگاسیان، هارتانیان، و گیروگیان^۳، ۲۰۱۴). نویسندگان مفهوم آستانه استراتژیک داده سری زمانی را به هر نوع اطلاعاتی که جریانی از اسناد و فعالیت‌ها هستند، گسترش می‌دهند. آنها مفهوم نرمال بودن آن جریان‌ها را در روش پیشنهادی خود پیاده‌سازی کردند و مکانیزمی را برای تشخیص ناهنجاری آنها در جریان زمان اجرا ایجاد کردند. تحت محدودیت‌های منحصربه‌فرد در پیچیدگی و مقیاس‌پذیری، آنها ساختار تصمیم‌گیری جدیدی را برای بازیابی اطلاعات از جریان داده‌ها پیاده‌سازی کردند. تکنیک پیشنهاد شده فوق، در درجه اول مبتنی بر تجزیه و تحلیل غیرعادی داده‌های رویداد از فایل‌های گزارش و به دست آوردن اطلاعات مفید از منبع و انواع رویدادها است.

الگوریتم‌های تشخیص تقلب بانکی

بسیاری از روش‌های ML برای مبارزه با تقلب توسعه داده شده است. تکنیک‌های متداول کشف تقلب عبارتند از ANN، ES (مبتنی بر دانش (KB))، موتور استنتاج و داده‌کاوی (ماراتون، ۲۰۱۳؛ وی و همکاران، ۲۰۱۳). بیشترین رویکردهای مورد استفاده برای کشف تقلب روش‌های نظارت شده، بدون نظارت، نیمه نظارت شده و ترکیبی هستند (بولتون و هند، ۲۰۲۰؛ جان و همکاران، ۲۰۱۶؛ عبدالله و همکاران، ۲۰۱۶؛ کانادول و بانجری^۴، ۲۰۰۹). کوچا و روجیرو^۵ (۲۰۱۱) یک سیستم کشف تقلب را برای بانکداری آنلاین با تمرکز بر تجزیه و تحلیل محلی و جهانی رفتار کاربران پیشنهاد

برای شناسایی موارد انحرافی است و شامل یک یا دو کار زیر است:

(الف) شناسایی داده‌های غیرعادی، به عنوان مثال، نویز، انحرافات یا نقاط پرت از مجموعه داده اصلی و (ب) کشف نمونه‌های داده جدید بر اساس دانش آموخته شده بر اساس مجموعه داده اصلی، تشخیص ناهنجاری را می‌توان برای اهداف مختلفی استفاده کرد (آگروال^۱، ۲۰۱۷). مانند تشخیص تقلب، تجزیه و تحلیل کیفیت داده‌ها، اسکن امنیتی، نظارت بر فرآیند و سیستم، نظارت تصویر/ویدئو، تشخیص هرزنامه، شناسایی حملات مخرب داخلی، پاک‌سازی داده‌ها قبل از آموزش مدل‌های آماری، تجزیه و تحلیل رفتار انسان (چیو، کیم، کانگ، ویوم^۲، ۲۰۲۱) و تشخیص خطای حسگر (درویشی، کیانزو، عید و راسی^۳، ۲۰۲۰). برخی از محققین، یک تکنیک تشخیص نفوذ مبتنی بر قانون (ID) پیشنهاد و اذعان می‌کنند که AD یکی از اولین رویکردها برای ID است (ایلگان، کیمر و پارس^۴، ۱۹۹۵). سایر محققین نشان دادند که سرقت حساب یا سرویس در رایانش ابری خطرناک‌تر است. نویسندگان چارچوبی با تکنیک AD برای نمایه کردن رفتارهای معمولی کاربر پیشنهاد کردند. هنگامی که یک نمایه کاربر از اطلاعات جمع‌آوری شده کشف می‌شود، هشدارها توسط همه رفتارهای مشکوک شناسایی شده توسط نمایه فعال می‌شوند. هشدارها به پایگاه داده ارسال می‌شود و به صاحب حساب و مدیر ابر اطلاع داده می‌شود. دو جزء اصلی چارچوب وجود دارد: بخش اول جمع‌آوری داده و بخش دوم مازول یادگیری است که تکنیک‌های مختلف ML را برای استخراج روندهای مکرر از داده‌های آموزشی و به دست آوردن محدودیت رتبه‌بندی مشکوک معرفی کرده است. اگر رتبه‌بندی مشکوک از حد تعیین شده بیشتر شود، تراکنش به عنوان یک ناهنجاری توسط طرح گزارش می‌شود چیو و همکاران^۵ (۲۰۱۳) پیشنهاد می‌کنند که

- 1- Aggarwal
- 2- Choi, Kim, Kang, & Youm
- 3- Darvishi, Ciunzo, Eide, & Rossi
- 4- Ilgun, Kemmerer, & Porras
- 5- Chiu
- 6- Chiu, Yeh, & Lee
- 7- Brabazon
- 8- Marvasti, Poghosyan, Harutyunyan, & Grigoryan
- 9- Chandola, & Banerjee
- 10- Kovach and Ruggiero

اين نتيجه رسيدند كه بالاترين خروجي در كشف قلب روي مجموعه داده تركيبي است. اندازه مجموعه داده مورد بحث قرار نمي گيرد، اما اين تكنيك براي مجموعه داده هاي بزرگ كه در آن مجموعه داده ديگري با تركيبي از مجموعه داده هاي اوليه مورد نياز است، بهينه نيست.

باي (۲۰۱۳) روشي را پيشنهاد مي كند كه در درجه اول يك گزينه جستجوي مؤثر براي اطلاعات بلادرنگ است، اما براي حوزه طبقه بندي مناسب نيست، چرا كه هر عمليات بايد به عنوان قلب يا غير قلب طبقه بندي شود. كواچ و روجيرو (۲۰۱۱) پيشنهاد مي كنند كه سيستم رديابي قلب بر رفتار مشتريان متمرکز است و به شدت به دسترسي دستگاه بستگي دارد كه براي اطلاعات بزرگ و در زمان واقعي مناسب نيست. راه حل كشف قلب كارت اعتباري پيشنهاد شده توسط دومان و اوزچليك (۲۰۱۱) از رويكردهاي فراابتكاري استفاده مي كند. با اين حال، نويسندگان روش هايي براي برخورد با اطلاعات و اطلاعات بزرگ در زمان واقعي ارائه نكرده اند.

وي و همكاران (۲۰۱۳) يك تكنيك كشف قلب را براي داده هاي بسيار نامتعادل با استفاده از الگوريتم Contrast Miner ارائه مي دهند كه الگوهاي تضاد را استخراج مي كند و قلب را از رفتار واقعي متمايز مي كند. كو و همكاران (۲۰۰۴) معتقدند كه تحقيقات كشف قلب عمدتاً از داده كاوي، آمار و هوش مصنوعي استفاده مي كند و قلب از ناهنجاري ها در داده ها و الگوها تشخيص داده مي شود.

چن و همكاران (۲۰۱۹) يك سيستم تشخيص قلب مبتني بر گراف را براي بيمه تجارت الكترونيكي به نام InfDetect مورد بحث قرار دادند. اين آزمون بر روي داده هاي شركت بيمه خصوصي از جمله بيمه سپرده تضميني و بيمه باربري برگشت انجام شد. مدل آنها از تركيبي از گراف ها و ويژگي هاي خام به عنوان داده ورودی، ساخت ويژگي از طريق مدل يادگيري گراف تحت نظارت و بدون نظارت به نام DeepWalk استفاده مي كند. رمزگذار خودكار حذف نويز و پردازش ويژگي نيز پياده سازي شده

مي كنند. ارزيايي افتراقي براي به دست آوردن شواهدی از قلب استفاده می شود که در آن تغییر قابل توجهی از رفتار عادی یک قلب بالقوه را نشان می دهد. شواهد قلب بر تعداد دسترسي هاي کاربر و مقدار احتمالي كه در طول دوره متفاوت است متمرکز می باشد. روش پيشنهادهی كشف قلب آنها بر شناسايي كارآمد دستگاه هاي مورد استفاده براي كنترل حساب ها و ارزيايي احتمال كلاه برداري با نظارت بر تعداد سوابق مجزايي كه هر دستگاه به آن دسترسي دارد، متمرکز دارد. روشي براي شناسايي قلب در كارت اعتباري توسط دومان و اوزچليك^۱ (۲۰۱۱) ارائه شده است. نويسندگان تركيبي از دوروش فراابتكاري معروف را براي رتبه بندي پيشنهاد كردند كه عبارتند از: الگوريتم هاي ژنتيك و جستجوي پراكنده. اين تكنيك بر روي داده هاي واقعي پياده سازي شده است و يافته هاي به دست آمده نسبت به سيستم فعلي در حال استفاده بسيار مؤثر هستند. هر تراكنش با اين رويكرد رتبه بندي می شود و عمليات بسته به نتيجه به صورت قلب يا مشروع طبقه بندي می شود. آنها معتقدند كه يك طرح FD بهتر از طرحي است كه بسياري از كلاه برداري هاي كم خطر را شناسايي می كند، كه جرم را حتي از نظر مقدار كمتر اما از نظر ارزش بيشتر كشف می كند. كو و همكاران^۲ (۲۰۰۴) نيز يك مطالعه مبتني بر داده كاوي از تحقيقات كشف قلب را براي دسته بندي تحقيقات مبتني بر چهار روش اصلي شامل رويكردهاي نظارت شده، تركيبي، نيمه نظارتی و بدون نظارت انجام دادند و همچنين رابطه كشف قلب را با ساير زمينه ها شناسايي كردند. يك استراتژي كشف قلب توسط هرلند، خوش گفتار، و بادر^۳ (۲۰۱۸) براي قلب مديكرو^۴ با استفاده از سه مجموعه داده خدمات مديكرو و مديكرو پيشنهاد شده است. روش آنها بر روي مجموعه داده تركيبي با اتصال چندين مجموعه داده آموزشي عمل می كند. نويسندگان از سه طبقه بندي كننده استفاده كردند: مدل هاي جنگل تصادفي، تقويت درخت گرايدان و مدل هاي رگرسيون لجستيك و از متريك ناحيه زير منحنی (ROC) براي اندازه گيري عملکرد FD استفاده كردند. آنها به

1- Duman and Ozcelik

2- Kou

3- Berland, Khoshgoftaar, and Bauder

4- Medicare

۳- يافته‌های تحقيق

رگرسیون لجستیک برای تشخیص تقلب

رگرسیون لجستیک^۳ به یک روش یادگیری نظارت شده اشاره دارد که با تصمیمات قطعی استفاده می‌شود. به این معنی که نتایج به دست آمده در صورت انجام معامله، تقلب یا غیر تقلب محسوب می‌شود. این رویکرد از یک رابطه علت و معلولی برای توسعه مجموعه داده‌های سازمان یافته استفاده می‌کند. تکنیک تحلیل رگرسیون زمانی که در تشخیص تقلب مورد استفاده قرار می‌گیرد به دلیل چندین متغیر و اندازه مجموعه داده پیچیده تر است. این مدل (الگوریتم) پیش بینی می‌کند که آیا تراکنش‌های جدید به عنوان تقلبی علامت گذاری می‌شوند یا خیر. این مدل‌ها معمولاً برای مشتریان خود از تجار بزرگ تر دقیق هستند، اما معمولاً مدل‌های عمومی همچنان قابل اجرا هستند.

درخت تصمیم برای تشخیص تقلب

الگوریتم‌های درخت تصمیم^۴ برای طبقه بندی فعالیت‌های غیر معمول در یک تراکنش از یک کاربر مجاز استفاده می‌شود. این الگوریتم‌ها دارای محدودیت‌هایی هستند که در مجموعه داده‌ها در طبقه بندی تقلب استفاده می‌شوند. الگوریتم‌های درخت تصمیم در طبقه بندی یا مشکلات مدل سازی برون یابی رگرسیون استفاده می‌شوند. آنها اساساً مجموعه قوانینی هستند که برای استفاده از موارد کلاه برداری شامل مشتریان مهارت دارند. ایجاد یک درخت تصمیم ویژگی‌های نامرتب را نادیده می‌گیرد و نیازی به عادی سازی داده‌های گسترده ندارد. هنگامی که یک درخت تحت بازرسی قرار می‌گیرد، دلیل اینکه چرا یک تصمیم خاص بسته به لیست قوانین فعال شده توسط یک مشتری خاص گرفته شده است، درک می‌شود. خروجی الگوریتم یادگیری ماشین ممکن است مدلی باشد که میمون درخت تصمیم است. این یک امتیاز احتمالی تقلب را بر اساس شرایط قبلی تعیین

است. آنها سپس یک درخت تصمیم گیری مبتنی بر گرادیان تقویت شده مبتنی بر سرور به نام PSMART را برای خروجی احتمال تقلب اعمال می‌کنند. محققان ادعا کرده اند که طرح پیشنهادی آنها به صرفه جویی میلیون ها دلار در سال برای این شرکت های تجارت الکترونیک کمک می‌کنند. یکی از ویژگی های کلیدی مهندسی شده توسط محققان، استفاده از اختصاص امتیازهای bin برای مبالغ تراکنش بود که به بهبود عملکرد کمک کرد.

آراجو و همکاران (۲۰۱۷) الگوریتم "BreachRadar" را توسعه داده اند، یک الگوریتم متناوب توزیع شده که احتمال در معرض خطر قرار گرفتن کارت بانکی را به مکان های مختلف ممکن برای استفاده از کارت اختصاص می‌دهد. این نقاط سازش^۱ (POC) نقطه شروع برای تقلب در تراکنش های بانکی است. یک مثال می تواند نقض داده ها باشد که منجر به کسب اطلاعات مشتریان بانکی می شود، از جمله اطلاعات کارت اعتباری آنها که سپس به صورت آنلاین از طریق سایت های DarkNet شسته می شود یا شخصاً توسط مجرمان استفاده می شود. مثال دیگر می تواند اسکیمینگ^۲ کارت باشد، که در آن جزئیات کارت مشتری با استفاده از دستگاهی که عمداً برای به دست آوردن غیرقانونی این اطلاعات تغییر داده شده، کپی یا شبیه سازی می شود. هدف از شناسایی POC جلوگیری از استفاده جعلی از داده های مشتری است. محققان مدلی با عملکرد بالا با دقت ۹۰ درصد و یادآوری در برابر مجموعه داده ای با ۱۰ درصد کارت های اعتباری مشتریان را ارائه کردند. این تکنیک شامل تشکیل یک مدل گراف دوبخشی برای نشان دادن همه کارت ها و مکان آنها است. با اجرای یک الگوریتم در حافظه که امکان به روزرسانی احتمالات POC را فراهم می‌کند، محققان اولین روش توزیع شده را تولید کرده اند که قادر به تشخیص خودکار POC ها می باشد.

1- Points-of-Compromise

2- Skimming

3- Logistic Regression

4- Decision Tree

می‌کند.

جنگل تصادفی برای تشخیص تقلب

جنگل تصادفی^۱ از ترکیبی از درختان تصمیم برای بهبود نتایج استفاده می‌کند. هر درخت معاملات را برای شرایط مختلف ارزیابی می‌کند (آیدیوارا، کی؛ آیدیوارا، وی، ۲۰۸). همچنین مجموعه داده‌های تصادفی آموزش داده می‌شوند. بسته به آموزش درخت تصمیم، هر درخت یک تراکنش را به عنوان تقلبی یا غیر متقلبانه طبقه‌بندی می‌کند. سپس از مدل برای پیش‌بینی نتیجه استفاده می‌شود. این به آشکارسازهای تقلب اجازه می‌دهد تا خطای موجود در درخت را یکسان کنند. دقت مدل عملکرد کلی را افزایش می‌دهد و در عین حال توانایی تفسیر نتایج را حفظ می‌کند و امتیازات قابل توضیحی را برای کاربران ما ارائه می‌دهد.

زمان‌های اجرا جنگل تصادفی سریع هستند و داده‌هایی را که گم شده یا نامتعادل هستند مدیریت می‌کنند. ML های جنگل تصادفی دارای نقاط ضعفی هستند، مثلاً هنگامی که در گرسیون استفاده می‌شوند، قادر به پیش‌بینی فراتر از تنوع در آموزش داده‌ها نیستند و ممکن است مجموعه‌های داده‌ای که نویز در نظر گرفته می‌شوند بیش از حد برآزش کنند.

شبکه‌های عصبی برای تشخیص تقلب

شبکه‌های عصبی^۲ مبتنی بر مغز انسان هستند. آنها از لایه‌های محاسباتی مختلفی استفاده می‌کنند. آنها از محاسبات شناختی استفاده می‌کنند که به ساخت ماشین‌هایی کمک می‌کند که می‌توانند از الگوریتم‌های خودآموزی که شامل داده‌کاوی، تشخیص الگوها و پردازش زبان طبیعی است استفاده کنند (گراپ، ۲۰۱۶). شبکه‌های عصبی برای فرآیند آموزش داده‌ها چندین لایه را پشت سر می‌گذارند. نتایج دقیق‌تری را در مقایسه با مدل‌های دیگر ارائه می‌دهد؛ زیرا از محاسبات شناختی استفاده می‌کند و از الگوهای رفتار مجاز می‌آموزد. بنابراین، بین معاملات «تقلبی» و غیر تقلبی تمایز قائل می‌شود. شبکه‌های

عصبی کاملاً سازگار هستند و از مجموعه الگوهای رفتار مشروع درس می‌گیرند. اینها با تغییر در رفتار آنچه که تراکنش‌های استاندارد تلقی می‌شوند سازگار می‌شوند و اشکال تراکنش‌های تقلب را شناسایی می‌کنند. فرآیند شبکه‌های عصبی سریع است و در زمان واقعی کار می‌کند.

۴- نتیجه‌گیری

فعالیت‌های متقلبانه در طول دوره کرونا افزایش یافته است و سازمان‌هایی با روش‌های سنتی کشف تقلب ثابت کرده‌اند که نتایج چشمگیری داشته‌اند؛ زیرا چشم انسان همیشه ناهنجاری‌ها را در سیستم بانکی تشخیص نمی‌دهد. کلاهبرداری مالی یکی از اصلی‌ترین مشکلاتی است که اعتماد مشتریان را تضعیف می‌کند و همچنین زیان اقتصادی به بانک‌ها و مؤسسات مالی وارد می‌کند. در سال‌های اخیر، هم‌زمان با گسترش تقلب، مؤسسات مالی به دنبال راه‌حلی برای یافتن راه‌حل مناسب در مبارزه با کلاهبرداری بودند. با توجه به تغییرات پیشرفته و متنوع در روش‌های تقلب، تحقیقات گسترده‌ای برای کشف تقلب انجام شده است. تقلب پیچیده بانکداری آنلاین نشان‌دهنده سوء استفاده یکپارچه از منابع در دنیای اجتماعی، سایبری و فیزیکی است. با این حال، اطلاعات بسیار محدودی برای تشخیص کلاهبرداری پویا از رفتار مشتری واقعی در چنین محیط داده‌ای بسیار کم و نامتعادل در دسترس است، که باعث می‌شود تشخیص فوری و مؤثر مهم و چالش برانگیز شود. بانک‌ها به دنبال کاهش زیان‌های هنگفت از طریق سیستم‌های تشخیص و پیشگیری از تقلب هستند. بسیاری از فن‌آوری‌های پیشرفته کلاهبرداری برای شناسایی و پیشگیری از تراکنش‌های بانکی اینترنتی تقلبی استفاده می‌شوند. با این حال، آنها هیچ مکانیسم شناسایی مؤثری برای شناسایی کاربران قانونی و ردیابی فعالیت‌های غیرقانونی آنها ندارند. توجه به نتایج به دست آمده از این مطالعه، توصیه می‌شود که از الگوریتم‌های یادگیری ماشین برای کشف تقلب استفاده شود تا جایگزین سیستم‌های تشخیص تقلب مبتنی بر قانون نادرست شود

1- Random Forest
2- Neural Networks

منابع:

- 5-Abdallah, A., Maarof, M. A., & Zainal, A. (2016). Fraud detection system: A survey. *Journal of Network and Computer Applications*, 68, 90-113.
- 6-ACI.(2011). Proactive Risk Manager. Retrieved from <https://www.aciworldwide.com/products/proactive-risk-manager>.
- 7-Aggarwal, C. C. (2017). Outlier Analysis. 2nd. Cham, Switzerland: Springer.
- 8-Amarasinghe, T., Aponso, A., & Krishnarajah, N. (2018, May). Critical analysis of machine learning based approaches for fraud detection in financial transactions. In *Proceedings of the 2018 International Conference on Machine Learning Technologies* (pp. 12-17).
- 9-APF. (2018). Western Australia Police Force. Retrieved from <https://www.police.wa.gov.au/>
- 10-Araujo, M., Almeida, M., Ferreira, J., Silva, L., & Bizarro, P. (2017, June). Breachradar: Automatic detection of points-of-compromise. In *Proceedings of the 2017 SIAM International Conference on Data Mining* (pp. 561-569). Society for Industrial and Applied Mathematics.
- 11-Ayyadevara, V. K., & Ayyadevara, V. K. (2018). Gradient boosting machine. *Pro machine learning algorithms: A hands-on approach to implementing algorithms in python and R*, 117-134.
- 12-Bao, Y., Hilary, G., & Ke, B. (2022). Artificial intelligence and fraud detection. *Innovative Technology at the Interface of Finance and Operations: Volume I*, 223-247.
- 13-Bilgin, M. H., MARCO LAU, C. K., & Demir, E. (2012). Technology transfer, finance channels, and SME performance: New evidence from developing countries. *The Singapore Economic Review*, 57(03),

1250020.

14-Brabazon, A., Cahill, J., Keenan, P., & Walsh, D. (2010, July). Identifying online credit card fraud using artificial immune systems. In *IEEE Congress on Evolutionary Computation* (pp. 1-7). IEEE.

15-Bolton, R. J., & Hand, D. J. (2002). Statistical fraud detection: A review. *Statistical science*, 17(3), 235-255.

16-Capterra. (2019). Fraud Management Software. Retrieved from <https://www.capterra.com.au/directory/10058/financial-fraud-detection/software>

17-Carminati, M., Polino, M., Continella, A., Lanzi, A., Maggi, F., & Zanero, S. (2018). Security evaluation of a banking fraud analysis system. *ACM Transactions on Privacy and Security (TOPS)*, 21(3), 1-31.

18-Carminati, M., Caron, R., Maggi, F., Epifani, I., & Zanero, S. (2015). BankSealer: A decision support system for online banking fraud analysis and investigation. *computers & security*, 53, 175-186

Bagga, S., Goyal, A., Gupta, N., & Goyal, A. (2020). Credit card fraud detection using pipeling and ensemble learning. *Procedia Computer Science*, 173, 104-112.

19-Chandola, V., & Banerjee, A. V., K. (2009). Anomaly detection: A survey. *ACM Computing survey*, 41.

20-Chen, C., Liang, C., Lin, J., Wang, L., Liu, Z., Yang, X., ... & Qi, Y. (2019, December). InfDetect: A large scale graph-based fraud detection system for E-commerce insurance. In *2019 IEEE International Conference on Big Data (Big Data)* (pp. 1765-1773). IEEE.

21-Chen, C. M., Guan, D. J., Huang, Y. Z., & Ou, Y. H. (2012, August). Attack sequence detection in cloud using hidden markov model. I

22-Chen, L., Zhang, Z., Liu, Q., Yang, L., Meng, Y., & Wang, P. (2019, May). A method for online transaction fraud detection based on individual behavior. In *Proceedings of the ACM Turing Celebration Conference-China* (pp. 1-8).

23-Chen, K., Yadav, A., Khan, A., & Zhu, K. (2020). Credit Fraud Detection Based on Hybrid Credit Scoring Model. *Procedia Computer Science*, 167, 2-8.

24-Chiu, C. Y., Yeh, C. T., & Lee, Y. J. (2013, December). Frequent pattern based user behavior anomaly detection for cloud system. In *2013 Conference on Technologies and Applications of Artificial Intelligence* (pp. 61-66). IEEE.

25-Choi, S., Kim, C., Kang, Y. S., & Youm, S. (2021). Human behavioral pattern analysis-based anomaly detection system in residential space. *The Journal of Supercomputing*, 77, 9248-9265.

26-CURTIS, C. V. (2012). The Fraud Triangle: Fraudulent Executives, Complicit Auditors and Intolerable Public Injury.

27-Darvishi, H., Ciuonzo, D., Eide, E. R., & Rossi, P. S. (2020). Sensor-fault detection, isolation and accommodation for digital twins via modular data-driven architecture. *IEEE Sensors Journal*, 21(4), 4827-4838.

28-Dehaven, V. R. (2014). Machine learning: Future capabilities and their implications.

29-Dong, F., Wang, H., Li, L., Guo, Y., Bissyandé, T. F., Liu, T., ... & Klein, J. (2018, October). Frauddroid: Automated ad fraud detection for android apps. In *Proceedings of the 2018 26th ACM Joint Meeting on European Software Engineering Conference and Symposium on the Foundations of Software Engineering* (pp. 257-268).

30-Duman, E., & Ozcelik, M. H. (2011). Detecting credit card fraud by genetic algorithm and scatter search. *Expert Systems with Applications*, 38(10), 13057-13063.

31-Eurostat, G. D. P., & Components, M. (2020). Available online: https://ec.europa.eu/eurostat/databrowser/view.NAMQ_10_GDP/default/table.

32-FBI. (2018). Internet Crime Complaint Center, 20182019(20/08/2019). Retrieved from https://pdf.ic3.gov/2018_I_C3_Report.Pdf

33-FICO. (2010). Fico Application Fraud Manager. Retrieved from <https://www.fico.com/en/products/fico-application-fraud-manager>

34-FraudNet. (1997). Enterprise Fraud Prevention. Retrieved from <https://fraud.net>

35-Gbudeanu, L., Brici, I., Mare, C., Mihai, I. C., & cheau, M. C. (2021). Privacy intrusiveness in

- financial-banking fraud detection. *Risks*, 9(6), 104.
- 36-Graupe, D. (2016). *Deep learning neural networks: Design and case studies*. World Scientific Publishing Company.
- 37-Hafiz, K. T., Aghili, S., & Zavorsky, P. (2016, June). The use of predictive analytics technology to detect credit card fraud in Canada. In *2016 11th Iberian Conference on Information Systems and Technologies (CISTI)* (pp. 1-6). IEEE.
- 38-Herland, M., Khoshgoftaar, T. M., & Bauder, R. A. (2018). Big data fraud detection using multiple medicare data sources. *Journal of Big Data*, 5(1), 1-21.
- 39-Herschel, G., Linden, A., & Kart, L. (2015). Magic quadrant for advanced analytics platforms. *Gartner Report G*, 270612.
- 40-Hopton, D. (2016). Proceeds of Crime Act 2002–Part 7: Requirements and Offences. In *Money Laundering* (pp. 53-78). Gower.
- 41-Horák, M., Stupka, V., & Husák, M. (2019, August). GDPR compliance in cybersecurity software: a case study of DPIA in information sharing platform. In *Proceedings of the 14th international conference on availability, reliability and security* (pp. 1-8).
- 42-Ilgun, K., Kemmerer, R. A., & Porras, P. A. (1995). State transition analysis: A rule-based intrusion detection approach. *IEEE transactions on software engineering*, 21(3), 181-199.
- 43-Jha, S., Guillen, M., & Westland, J. C. (2012). Employing transaction aggregation strategy to detect credit card fraud. *Expert systems with applications*, 39(16), 12650-12657.
- 44-John, S. N., Anele, C., Kennedy, O. O., Olajide, F., & Kennedy, C. G. (2016, December). Realtime fraud detection in the banking sector using data mining techniques/algorithm. In *2016 international conference on computational science and computational intelligence (CSCI)* (pp. 1186-1191). IEEE.
- 45-Kou, Y., Lu, C. T., Sirwongwattana, S., & Huang, Y. P. (2004, March). Survey of fraud detection techniques. In *IEEE International Conference on Networking, Sensing and Control, 2004* (Vol. 2, pp. 749-754). IEEE.
- 46-Kovach, S., & Ruggiero, W. V. (2011, February). Online banking fraud detection based on local and global behavior. In *Proc. of the Fifth International Conference on Digital Society, Guadeloupe, France* (pp. 166-171).
- 47-Lavion, D. (2018). Global Economic Crime and Fraud Survey 2018, 30. Retrieved from Pulling fraud out of the shadows website: <https://www.pwc.com/gx/en/forensics/global-economic-crime-and-fraud-survey->
- 48-Lee, J. H., Park, M. W., Eom, J. H., & Chung, T. M. (2011, February). Multi-level intrusion detection system and log management in cloud computing. In *13th International Conference on Advanced Communication Technology (ICACT2011)* (pp. 552-555). IEEE.
- 49-Legg, C. (2007). Ontologies on the semantic web. *Annual review of information science and technology*, 41, 407.
- 50-Li, Z., Huang, M., Liu, G., & Jiang, C. (2021). A hybrid method with dynamic weighted entropy for handling the problem of class imbalance with overlap in credit card fraud detection. *Expert Systems with Applications*, 175, 114750.
- 51-Marican, L., & Lim, S. (2014). Microsoft Consumer Safety Index reveals impact of poor online safety behaviours in Singapore. Retrieved from <https://news.microsoft.com/en-sg/2014/02/11/microsoft-consumer-safety-index-reveals-impact-of-poor-online-safety>.
- 52-Maruatona, O. (2013). *Internet banking fraud detection using prudent analysis* (Doctoral dissertation, University of Ballarat).
- 53-Maruti T. (2021). How Machine Learning Facilitates Fraud Detection?. [online] Available at:<<https://marutitech.com/machine-learning-fraud->
- 54-Marvasti, M. A., Poghosyan, A. V., Harutyunyan, A. N., & Grigoryan, N. M. (2014). An enterprise dynamic thresholding system. In *11th International Conference on Autonomic Computing ({ICAC} 14)* (pp. 129-135).

- 55-Misra, S., Thakur, S., Ghosh, M., & Saha, S. K. (2020). An autoencoder based model for detecting fraudulent credit card transaction. *Procedia Computer Science*, 167, 254-262.
- 56-Moon, W. Y., & Kim, S. D. (2017). Adaptive fraud detection framework for fintech based on machine learning. *Advanced Science Letters*, 23(10), 10167-10171.
- 57-Nanduri, J., Jia, Y., Oka, A., Beaver, J., & Liu, Y. W. (2020). Microsoft uses machine learning and optimization to reduce e-commerce fraud. *INFORMS Journal on Applied Analytics*, 50(1), 64-79.
- 58-Nathan, R. J., Victor, V., Gan, C. L., & Kot, S. (2019). Electronic commerce for home-based businesses in emerging and developed economy. *Eurasian Business Review*, 9, 463-483.
- 59-Nicholls, J., Kuppa, A., & Le-Khac, N. A. (2021). Financial cybercrime: A comprehensive survey of deep learning approaches to tackle the evolving financial crime landscape. *Ieee Access*, 9, 163965-163986.
- 60-Orek, M., Örek, E., & Bahtiyar, . (2019, May). A deep learning method for fraud detection in financial systems: Poster. In *Proceedings of the 12th Conference on Security and Privacy in Wireless and Mobile Networks* (pp. 298-299).
- 61-Olowookere, T. A., & Adewale, O. S. (2020). A framework for detecting credit card fraud with cost-sensitive meta-learning ensemble approach. *Scientific African*, 8, e00464.
- 62-Ocansey, E. O. N. D., & Ganu, J. (2017). The role of corporate culture in managing occupational fraud. *Research Journal of Finance and Accounting*, 8(24).
- 63-Pascual, A., Marchini, K., & Miller, S. (2017). Identity fraud: securing the connected life. *Retrieved on January, 20, 2018*.
- 64-PattemSpy. (2015). PattemSpy For Banking - Fraud Management Software. Retrieved from <https://www.pattemspy.tech/>
- 65-Pinto, S. O., & Sobreiro, V. A. (2022). Literature review: Anomaly detection approaches on digital business financial systems. *Digital Business*, 100038.
- 66-Politou, E., Alepis, E., & Patsakis, C. (2019). Profiling tax and financial behaviour with big data under the GDPR. *Computer law & security review*, 35(3), 306-329.
- 67-Phua, C., Lee, V., Smith, K., & Gayler, R. (2010). A comprehensive survey of data mining-based fraud detection research. *arXivpreprintarXiv:1009.6119*.
- 68-Richards, D. (2003). Knowledge-based system explanation: The ripple-down rules alternative. *Knowledge and Information Systems*, 5(1), 2-25.
- 69-RiskNet. (1998). Fraud Managed Services. Retrieved from <https://www.aicorporation.com/products-services/fraud-managed-services/>
- 70-SAS. (2007). SAS Fraud Management. Retrieved from https://www.sas.com/en_au/software/fraud-management.html
- 71-Sando, S., 2021. Consumer Preference Drives Shift in Authentication. [online] Javelin. Available at: <https://www.javelinstrategy.com/coverage-area/>
- 72-cheau, M. C., Gafta, V. N., Achim, M. V., & Cotoc, C. N. B. (2020, October). Cyber Security Reactivity in Crisis Times and Critical Infrastructures. In *2020 24th International Conference on System Theory, Control and Computing (ICSTCC)* (pp. 691-698). IEEE.
- 73-Shihembetsa, E. (2021). *Use of artificial intelligence algorithms to enhance fraud detection in the Banking Industry* (Doctoral dissertation, University of Nairobi).
- 74-Stalla-Bourdillon, S., Pearce, H., & Tsakalakis, N. (2018). The GDPR: A game changer for electronic identification schemes? The case study of Gov. UK Verify. *Computer Law & Security Review*, 34(4), 784-805.
- 75-Stojanov, M. (2019). Protection Against Fraud In Electronic Trade Payments. 21, 9(1-ENG), 48-66.
- 76-Sudharsan, K., Kumar, V. A., Venkatesan, R., Sathyapreiya, V., & Saranya, G. (2019). Two three step authentication in ATM machine to transfer money and for voting application. *Procedia Computer Science*, 165, 300-306.
- 77-Syed, N. A., Liu, H., & Sung, K. K. (1999, August). Handling concept drifts in incremental learning

- with support vector machines. In *Proceedings of the fifth ACM SIGKDD international conference on Knowledge discovery and data mining* (pp. 317-321).
- 78-Wang, Y., & Wang, L. (2019, May). Bot-like Behavior Detection in Online Banking. In *Proceedings of the 4th International Conference on Big Data and Computing* (pp. 140-144).
- 79-Wei, W., Li, J., Cao, L., Ou, Y., & Chen, J. (2013). Effective detection of sophisticated online banking fraud on extremely imbalanced data. *World Wide Web*, 16, 449-475.
- 80-Williams, M. (2022). Elizabeth Holmes and Theranos: A play on more than just ethical failures. *Business Information Review*, 39(1), 23-31.

©Authors, Published by Journal of Intelligent Knowledge Exploration and Processing. This is an open-access paper distributed under the CC BY (license <http://creativecommons.org/licenses/by/4.0/>).



مقاله پژوهشی

مرکزیت در شبکه‌های اجتماعی مبتنی بر مدل آستانه خطی

Doi: 10.30508/KDIP.2022.361686.1049

علی سرآبادانی^۱ | خیرالله رهسپارفرد (نویسنده مسئول)^۲

۱- دانشجوی دکتری تخصصی مهندسی فناوری اطلاعات گرایش تجارت الکترونیک، دانشگاه قم، قم، ایران

۲- استادیار گروه مهندسی کامپیوتر و فناوری اطلاعات، دانشگاه قم، قم، ایران

تاریخ دریافت: ۱۴۰۱/۰۶/۲۰

تاریخ پذیرش: ۱۴۰۱/۱۰/۰۳

صفحه: ۰۰ - ۰۰

چکیده:

مرکزیت و انتشار نفوذ دو مؤلفه اساسی در تحلیل شبکه‌های اجتماعی هستند. در سال‌های اخیر، اندازه‌گیری مرکزیت توجه بسیاری از محققان را جلب کرده و همین امر باعث شده که بسیاری از مطالعات جدید در مورد تحلیل شبکه‌های اجتماعی و کاربردهای آن صورت بگیرد. مدل‌های متداول قدیمی از انتشار نفوذ به صورت گسترده برای اندازه‌گیری مرکزیت استفاده نمی‌کنند. بسیاری از مطالعات صورت گرفته بر اساس مدل آبشاری مستقل شکل گرفته‌اند. در این مقاله به بررسی ارزیابی و اندازه‌گیری مرکزیت بر اساس مدل آستانه خطی در شبکه‌های اجتماعی برچسب دار و وزن دار پرداخته شده است. یک روش برای اندازه‌گیری مرکزیت به منظور رتبه‌بندی اعضای شبکه در نظر گرفته شده است که رتبه آستانه خطی نامگذاری شده است. برای ارزیابی روش پیشنهادی چهار شبکه اجتماعی را با دو معیار اندازه‌گیری مرکزیت و یک روش آبشاری مستقل مقایسه شده‌اند. نتیجه نشان می‌دهد که روش پیشنهادی برای رتبه‌بندی اعضای شبکه و شبکه‌ها مناسب‌تر است. یک سنجه مرکزیت جدید مبتنی بر مدل آستانه خطی که برای هر عضو، به عنوان تعداد اعضای در نظر گرفته می‌شد که می‌توانست انتشار نفوذ صورت دهد، ارائه شد. این سنجه را با سه سنجه مرکزیت دیگر (مرکزیت کاتز، رتبه صفحه و رتبه‌بندی آبشاری مستقل) بر اساس نفوذ و تأثیر در چهار شبکه واقعی دو شبکه بزرگ (یکی جهت‌دار، دیگری بی‌جهت) و دو شبکه کوچک (یکی جهت‌دار، دیگری بی‌جهت) مقایسه شده است.

کلمات کلیدی: مرکزیت، مدل آستانه خطی، مدل آبشاری مستقل، گسترش نفوذ، شبکه اجتماعی

۱- مقدمه

مرکزیت یکی از مؤلفه‌های اساسی در تحلیل شبکه‌های اجتماعی است. بررسی این مؤلفه از سال ۱۹۴۸ آغاز شده است (باولاس^۱، ۱۹۴۸). یک شبکه اجتماعی را می‌توان مانند یک گراف نشان داد که گره‌ها نماینده اعضای شبکه هستند و یال‌ها نماینده روابط بین اعضا (سان و تانگ^۲، ۲۰۱۱؛ سدیمان^۳، ۱۹۸۳). گاهی یال‌ها وزن دار هستند که نشان‌دهنده قدرت بین هر کدام از روابط اعضا خواهند بود. در این میان سنجه‌های مرکزیت نشان‌دهنده ارتباط ساختاری هر عضو با این شبکه اجتماعی هستند. یکی از قدیمی‌ترین سنجه‌های مرکزیت، می‌توان به مرکزیت درجه^۴، مرکزیت بینایی^۵، مرکزیت نزدیکی^۶ اشاره کرد که به توپولوژی گراف نیز مرتبط بودند. در این سنجه‌ها، یک عضو بیشتر مرکزی در نظر گرفته می‌شود هنگامی که از درجه بالاتری برخوردار باشد، یا به دیگر اعضا نزدیک‌تر باشد، یا در بین کوتاه‌ترین مسیر بین اعضا قرار گرفته باشد (بالابس^۷، ۱۹۸۴؛ فریمن^۸، ۲۰۰۲). در سال‌های اخیر به دلیل افزایش حیرت‌برانگیز کاربران اینترنت و ایجاد شبکه‌های گوناگون و پیچیده اجتماعی، ظهور سنجه‌های مرکزیت پیچیده‌تر را به وضوح می‌توان مشاهده کرد (بادر و هوگو^۹، ۲۰۰۳؛ لوزارس، لپزولدان، بولیبار، و مانتانیول^{۱۰}، ۲۰۱۵). بنا بر حجم بسیار زیاد این شبکه‌ها، یعنی تعداد بسیار زیاد اعضا

و روابط بین آن‌ها، لازم می‌باشد که سنجه‌ها به صورت عملی قابل محاسبه باشند. امروزه، سنجه‌هایی وجود دارند که بر اساس این که چه مقدار اطلاعات از گره در شبکه پخش می‌شود یا به کار گرفته می‌شوند (الطاف امین و همکاران^{۱۱}، ۲۰۰۳؛ کیتساک و همکاران^{۱۲}، ۲۰۱۰؛ گومز، فیگوریا، و یوسیبو^{۱۳}، ۲۰۱۳). همچنین شاخص‌هایی برای شبکه‌های اجتماعی به خصوصی ایجاد شدند مانند شبکه تویتر که بیش از هفتاد سنجه مرکزیت فقط از سال ۲۰۱۰ برای آن تعیین شده‌اند (ریکولم، گونزال، مولینرو و ساما^{۱۴}، ۲۰۱۸). دو تا از شناخته شده‌ترین سنجه‌ها، سنجه رتبه صفحه (پیج، برین، موتاوانی، و وینگورال^{۱۵}، ۱۹۹۹)، و مرکزیت کاتز (کاتز^{۱۶}، ۱۹۵۳) هستند. هر دوی این سنجه‌ها، مشتق شده از مرکزیت بردار ویژه^{۱۷} هستند (بوناریچ^{۱۸}، ۱۹۷۲). از آنجایی که ارتباط بین کاربران برای بسیار از کاربری‌ها و سامانه‌ها کاربرد دارد مانند بازاریابی و پیروسی (دامیونوگ و ریچاردسون^{۱۹}، ۲۰۰۱)، انتشار اطلاعات (گاتلر، و پاترینگانی^{۲۰}، ۲۰۰۴؛ گراهل، گاه، لیل نیل، و تامکینز^{۲۱}، ۲۰۰۴)، استراتژی‌های جست‌وجو (آدامیک، و آدر^{۲۲}، ۲۰۰۵)، پیشنهاددهی حرفه‌ای و تخصصی (سانگ، تسنگ، لین، و سان^{۲۳}، ۲۰۰۶)، دستگاه‌های مجتمع (ژانگ، ژوو، وانگ، و ژائو^{۲۴}، ۲۰۱۳)، مدیریت ارتباط مشتری اجتماعی تئوری نفوذپذیری (مورون و میکس^{۲۵}، ۲۰۱۵)، فراتر از این موارد سنجه‌های مرکزیت برای شناسایی فعال‌ترین،

- 1- Bavelas
- 2- Sun, & Tang
- 3- Seidman
- 4- Centrality
- 5- Betweenness
- 6- Closeness centrality
- 7- Bollobás
- 8- Freeman
- 9- Bader, & Hogue
- 10- Lozares, López-Roldán, Bolibar, & Muntanyola
- 11- Altaf-Ul-Amine, etal
- 12- Gómez, Figueira, & Eusébio
- 13- Riquelme, Gonzalez-Cantergiani, Molinero, & Serna
- 14- Page, Brin, Motwani, & Winograd
- 15- Katz
- 16- Eigen Vector
- 17- Bonacich
- 18- Domingos, & Richardson
- 19- Gaertler, & Patrignani
- 20- Gruhl, Guha, Liben-Nowell, & Tomkins
- 21- Adamic, L., & Adar, E. (2005)
- 22- Song, Tseng, Lin, & Sun
- 23- Zhang, Zhu, Wang, & Zhao
- 24- Morone & Makse

و سنجه‌های مرکزی سازی ارائه شده را برای ۴ شبکه اجتماعی در نظر گرفته شده است. دوتا از آن‌ها شبکه‌های بزرگی محسوب می‌شوند: شبکه هیگز (جهت دار) و شبکه آرکایو (غیر جهت دار). دو شبکه دیگر شبکه‌های کوچک ولی مشهوری هستند. شبکه میز ناهارخوری (جهت دار)، شبکه اجتماعی دلفین‌ها (غیر جهت دار). هر سنجه مرکزیت یک شاخص مرکزیت جداگانه‌ای را ارائه می‌دهد مگر برای شبکه اجتماعی دلفین‌ها، جایی که روش مقام آستانه خطی جدید^{۱۳}، رتبه صفحه و مرکزیت کاتز شبیه به هم عمل می‌کنند.

۲- مبانی نظری

در این بخش به معرفی سنجه‌های مرکزیت شناخته شده پرداخته شده است. همچنین بررسی شده است که هر کدام به چه صورت کار می‌کنند و چگونه سنجه‌های مرکزیت همبستگی دارند. در نهایت سنجه‌های مرکزی سازی را بیان شده است. شبکه اجتماعی مانند یک گراف می‌باشد که در آن V نماینده گره‌ها و E نماینده یال‌های ارتباطی خواهند بود. بعضی اوقات نیاز به یک گراف وزن دار است که در آن G گراف مورد نظر و W تابع وزنی است که به هر گره تخصیص داده می‌شود.

معیارهای مرتبط با مرکزیت: یکی از پرکاربردترین سنجه‌ها در بین شاخص‌های مربوط، مرکزیت بردار ویژه هست (بانریچ و همکاران، ۱۹۷۲)، که مشخص می‌کند یک عضو در یک گروه مهم می‌باشد اگر همسایه‌های مهمی به آن لینک شده باشند یا اعضای خیلی زیادی با آن همسایه باشد. به بیان دیگر، یک ماتریس مجاورت A را در نظر بگیرید، عنصر a_{ij} از ماتریس A ، مقدار یک را دارد اگر عضو i گره j به عضو متصل شده باشد در غیر این صورت مقدار

محبوب‌ترین و تأثیرگذارترین کاربر در یک شبکه نیز استفاده می‌شوند (لی و همکاران، ۲۰۱۴). مدل انتشار تأثیر مدلی است که در آن اعضا بر روی یکدیگر توسط ارتباط بین خود با سایر اعضا در شبکه تأثیر می‌گذارند. گره‌ها تأثیر خود را در گراف اعمال می‌کنند. به این صورت که اگر مجموعه‌ای از اعضا به یک‌روند مطابق شوند، ممکن خواهد بود بر روی دیگر اعضا تأثیر بگذارند که آنان نیز به این روند تطبیق پیدا کنند. مطمئناً این یکی واضح‌ترین و شناخته‌ترین پدیده در تحلیل شبکه‌های اجتماعی است (یاسلی، و کلینبرگ، ۲۰۱۰). یکی از شناخته‌ترین مدل‌های عمومی برای انتشار نفوذ، مدل آستانه خطی است (وانگ، ما، وانگ، ایکسی، کوه، و چوانگ، ۲۰۲۱؛ کمپ، کلینبرگ، و تاردوس، ۲۰۰۵). مدل آستانه خطی بر اساس یک سری ایده از رفتار جمعی الگو گرفته (گرانووتر، ۱۹۷۸؛ اسچلینگ، ۲۰۰۶)، مدل آبشاری مستقل برای موضوع بازاریابی مطرح شده بود (گلدنبرگ، لیبای و مولر، ۲۰۰۱)، بسیاری از تلاش‌های تحقیقاتی به مطالعه مسئله پیشینه‌سازی تأثیر بر اساس مدل آستانه خطی و دیگر مدل‌ها اختصاص داده شده است (کندال، ۱۹۸۳؛ لیو و همکاران، ۲۰۱۶؛ فاستر، موت، پوترات، و روتنبرگ، ۲۰۰۱). در این مسئله تلاش برای پیدا کردن تعداد k عضو کلیدی برای پیشینه‌سازی انتشار تأثیر در میان تمام مجموعه‌ها با سایز یکسان است. بنابراین مسئله پیشینه‌سازی تأثیر بر اساس مدل آستانه خطی یک مسئله NP-hard خواهد بود (سارامکی، کیولا، اونلا، کاسکی، و کرتز، ۲۰۰۷؛ کورو متیوز، ۲۰۲۰؛ لانگیول، مییر، ۲۰۰۴) مطالعات در مورد سنجه‌های مرکزیت مشتق شده ناقص بوده و فقط مدل آبشاری مستقل را در نظر گرفته بودند این رتبه باندی‌ها در نظر گرفته و محاسبه شدند تا یک سری راه حل برای مسئله پیشینه‌سازی تأثیر ارائه شود. مرکزیت

- 1- Li et al
- 2- Easley, & Kleinberg
- 3- Wang, Ma, Wang, Xie, Koh, & Cheong
- 4- Kempe, Kleinberg, & Tardos
- 5- Granovetter
- 6- Schelling
- 7- Goldenberg, Libai, & Muller
- 8- Kendall
- 9- Foster, Muth, Potterat, & Rothenberg,
- 10- Saramäki, Kivelä, Onnela, Kaski, & Kertesz
- 11- Corr, & Matthews
- 12- Langville, & Meyer
- 13- LTR

مرکزیت مرتبط با نفوذ: در کنار تمام این سنجه‌هایی که بر اساس مرکزیت بردار ویژه بنا شده‌اند، یک سری سنجه‌های مرکزیت وجود دارند که بر اساس نفوذ تأثیر مبتنی بر مدل آبخاری مستقل بنا شده‌اند؛ بنابراین مدل، هر عضوی احتمال دارد که بر عضوی که هدف قرار گرفته تأثیرگذار باشد. برای انتشار این تأثیر، عضو باید فعال باشد و فقط یک شانس برای تأثیر بر عضوی دیگر دارد. وقتی یک عضو سعی در تأثیر بر عضوی دیگر را دارد این عضو فعال در نظر گرفته می‌شود و این روال برای این عضو تکرار می‌شود. تمام این فرآیند زمانی خاتمه می‌یابد که نه عضو فعالی وجود داشته باشد نه شانس برای تأثیرگذاری. به عنوان مثال، سنجه زیر را در نظر بگیرید (کمپ و همکاران، ۲۰۰۵)، که به آن رتبه‌بندی آبخاری^۳ مستقل گفته می‌شود.

$$ICR(u, p) = \frac{|F'(u, p)|}{\max \{|F'(v, p)| \mid v \in V(G)\}}$$

که در آن V مجموعه تمام اعضا در شبکه هست و $F'(u, p)$ فرآیند تأثیر نفوذ مبنی بر مدل آبخاری مستقل خواهد بود؛ که با عضوی با نشان u آغاز خواهد شد. این سنجه یک احتمال ثابت p برای تأثیر هر عضو در نظر می‌گیرد. بنابراین هر عضو به احتمال $p-1$ به صورت غیرفعال برای تأثیر از همسایه خود باقی مهم‌اند؛ و به احتمال $(p-1)$ به توان r از تعداد r همسایه که به آن متصل هستند غیرفعال باقی می‌ماند. در همه سنجه‌های مبتنی بر مدل آبخاری مستقل، مانند رتبه‌بندی آبخاری مستقل، نودهای فعال به احتمال انتشار بستگی دارند بنابراین یک سنجه می‌تواند رتبه‌بندی‌های متفاوتی برای هر اجرا داشته باشد. **مقایسه سنجه‌های مرکزیت:** به منظور مقایسه نتایج حاصل از دو سنجه مرکزیت، به طور معمول از همبستگی ایستا استفاده می‌شود. پرکاربردترین ضرایب برای ایجاد همبستگی در سنجه‌های مرکزیت، ضریب همبستگی رتبه‌ای اسپیرمن ρ و ضریب همبستگی رتبه‌ای کندال T

صفر را خواهد داشت. مرکزیت بردار ویژه از عضو u که با $EV(u)$ نشان داده می‌شود عبارت است از:

$$EV(u) = \frac{1}{\lambda} \sum_{v \in V(G)} (\alpha_{uv}) EV(v)$$

که در آن ثابت λ اندا، مقدار ویژه خوانده می‌شود. مرکزیت بردار ویژه زمانی نتایج قابل قبولی ارائه می‌دهد که گراف به شدت متصل باشد، مانند شبکه غیر جهت‌دار با کامپوننت^۲ های بسیار متصل. در شبکه‌های جهت‌دار واقعی، نودهایی به دست آورده می‌شود که مرکزیت بردار ویژه آن‌ها برابر صفر خواهد بود، بنابراین این سنجه ناکارآمد می‌شود. برای مثال، نودهایی که به کامپوننت به شدت متصل، نزدیک هستند ولی راهی برای اتصال به آن‌ها نیست. اگرچه، مرکزیت کاتز (کاتز، ۱۹۵۳) (براین کمبود بردار ویژه غلبه می‌کند ولی با استفاده از تخصیص دادن یک مقدار کوچک از مرکزیت برای نقش اعضای آزاد و بی‌اهمیت در شبکه مرکزیت کاتز یک عضو u که با $Katz(u)$ نمایش داده می‌شود عبارت است از:

$$KATZ(u) = \alpha \sum_{v \in V(G)} (\alpha_{vu}) KATZ(v) + \beta$$

که β یک ثابت که مستقل از شبکه است و α که فاکتور

محرك نامیده می‌شود، یک عدد بین صفر و $\frac{1}{\max}$ که اندای ماکس بزرگ‌ترین مقدار ویژه از A خواهد بود. لازم به ذکر است که وقتی α بیشترین مقدار و β برابر صفر باشد مقدار مرکزیت کاتز برابر مرکزیت بردار ویژه خواهد شد. علاوه بر این، یک سنجه مرکزیت شناخته شده دیگر نیز وجود دارد به نام رتبه صفحه (پیچ، ۱۹۹۹) که با $pr(u)$ نشان داده می‌شود.

$$PR(u) = (1 - \alpha) + \alpha \sum_{v \in V(G)} \frac{(\alpha_{vu}) PR(v)}{\delta + (v)}$$

1- Eigen Value
2- Component
3- ICR

$$C_i = \frac{2T(i)}{\delta(i)(\delta(i)-1)}$$

$\delta(i)$ درجه گره i هست. برای گراف های وزن دار، این مقادیر متفاوت خواهند بود. ولی کتابخانه NetworkX از رابطه زیر استفاده می کند.

$$C_i = \frac{1}{\delta(i)(\delta(i)-1)} \sum_{\substack{j,k \in V(G) \\ j \neq k}} (\hat{W}_{ij} \hat{W}_{ik} \hat{W}_{jk})^{1/3}$$

مدل آستانه خطی برای انتشار تأثیر

یک شبکه ی اجتماعی در قالب یک گراف قابل نمایش است که در آن گره ها اعضای شبکه و یال ها روابط فردی بین اعضا را نمایش می دهد. گاهی اوقات یال ها دارای وزن نیز هستند که این وزن بیان گر قوت هر رابطه است. یک ویژگی کمتر رایج، برچسب گره ها است که برای نمایش ویژگی های مختلف شبکه استفاده می شود. در این کار، برچسب ها بیان گر میزان مقاومت گره در تأثیرپذیری است، در حالی که وزن یال ها بیان گر قدرت تأثیر اعمال شده از یک عضو به عضو دیگر هست. یک شبکه ی اجتماعی وزن دار و برچسب دار، به صورت رسمی تر در قالب یک گراف تأثیر (مولینرو و همکاران، ۲۰۱۵) قابل تعریف است. در ادامه، در تمام موارد مدل آستانه ی خطی را برای تأثیر انتشار در نظر می گیریم.

تعریف ۱: یک گراف تأثیر یک چندتایی (G, w, f) است که در آن یک $G = (V, E)$ گراف جهت دار تشکیل شده از مجموعه ی افراد V و مجموعه ی یال های جهت دار E است. $E; w: E \rightarrow \mathbb{N}$ یک تابع وزن دهی است که یک وزن به هر یال تخصیص می دهد. $f: V \rightarrow \mathbb{N}$ یک تابع برچسب زنی است که میزان سهولت تأثیرپذیری هر عضو را اندازه گیری می کند. یک عضو $i \in V$ تنها در صورتی روی عضو $j \in V$ تأثیر می گذارد که $(i, j) \in E$.

توجه داشته باشید که می توان یک گراف بدون جهت را به صورت یک گراف جهت دار متقارن در نظر گرفت، بنابراین این معیار برای این نوع گراف نیز قابل استفاده است. با توجه به گراف (G, w, f) و یک مجموعه ی فعال ساز $X \subseteq V$ ، فرآیند تکراری فعال سازی زیر را در نظر

هستند. اگر w و γ دو لیست از n عضو باشند ما داریم (به و و، ۲۰۱۳).

$$\rho = 1 - \frac{6 \sum_{i=1}^n (x_i - y_i)^2}{n(n^2 - 1)} \text{ and } \tau = \frac{n_c - n_d}{0.5n(n-1)}$$

که x_i و y_i رتبه های اعضای i در لیست x و y هستند. علاوه بر این n_c تعداد زوج های موافق است و n_d تعداد زوج های ناموافق. مقدار هر دو ρ و τ در بازه -1 تا 1 هستند، به طوری که یک معنی این است که دو سنج برابر هستند و صفر یعنی این که دو تا کاملاً مستقل هستند و -1 به این معنا که دو خلاف یکدیگر هستند. محاسبه ضرایب همبستگی یک همگونی با مقدار احتمال خواهد داشت. به طور معمول، ما ۵۰٪ را به عنوان خط برش مشخص در نظر می گیریم بنابراین فرضیه در صورتی که احتمال کمتر از ۵۰٪ باشد رد خواهد شد.

سنجه های مرکزی سازی: برای هر سنج مرکزیت، سنج مرکزی سازی فریمن نیاز به دانستن بیشترین مقدار ممکن جمع تمام مرکزیت ها را برای تمامی شبکه ها دارد. این مورد برای بسیاری از سنج های مرکزیت مانند مرکزیت درجه، بینابینی و نزدیکی در شبکه های غیر جهت دار آسان خواهد بود ولی در مورد دیگر سنج ها، این مورد سخت خواهد بود. علاوه بر این، ما می توانیم ضریب خوشه بندی متوسط را به عنوان سنج مرکزی سازی بیان کنیم (واتس و استروگاتز، ۱۹۹۸). این ضریب، متوسط تمامی ضرایب خوشه بندی محلی است ولی در شبکه های وزن دار تعریف دیگری دارد. این مورد با رابطه زیر معرفی می شود.

$$ACC = \frac{1}{n} \sum_{i=1}^n C_i$$

C_i ضریب خوشه بندی محلی برای عضو i خواهد بود. برای گراف های غیر وزن دار برای مثال وقتی که تمام یال های ارتباطی وزنی برابر یک داشته باشند C_i برابر خواهد بود با تعداد مثلث هایی که i در آن عضو هست در برابر تمام تعداد مثلث های ممکن

where $\text{neighbors}(i) = \{i \in V | (i, j) \in E \vee (j, i) \in E\}$

در تعریف ما G یک گراف جهت دار است اما همسایه‌های عامل i عواملی هستند که به وسیله‌ی لبه در هر جهت به متصل شده‌اند. این به ما اجازه می‌دهد که فعال سازی اولیه را افزایش دهیم. مشاهده کنید که آن عوامل با زاویه‌ی خروجی کوچک نمی‌توانند نفوذ خود را در شبکه انتشار دهند. بدین ترتیب معیار به دست آمده شبیه به درجه‌ی مرکزیت است. از این رو، ما اجتماع $F_0(X) = \{i\}$ و همسایه‌های i را به عنوان فعال سازی اولیه در نظر می‌گیریم. هنگامی که $F(X)$ می‌تواند در زمان مشخص پردازش محاسبه شود، $LTR(i)$ نیز قابل محاسبه در این زمان است. علاوه بر این، به عنوان اندازه‌گیری رتبه‌بندی آشناری و به جای آن رتبه صفحه و مرکزیت کاتز، مقام آستانه خطی جدید همیشه همگراست. این نکته قابل توجه است که ما همچنین می‌توانیم معیارهای دیگر مربوط به همسایه‌های عامل را در نظر بگیریم. همچنین این مجموعه‌ها می‌توانند با توجه به وزن‌ها به منظور تنظیم نقش عامل مرکزی در گسترش نفوذ پالایه شوند. به عنوان مثال ما در فعال سازی اولیه می‌توانیم تنها آن دسته از همسایگان را که برای w_{ij} مطابق معیار پایین قرار دارند را در نظر بگیریم. تحت تأثیر دوم، نفوذ گسترش $F(\{i\})$ عامل i می‌تواند با اندازه‌گیری مقام آستانه خطی جدید همبستگی بیشتری داشته باشد. با این حال اگر محدودیت‌های همسایگان خیلی زیاد است، اندازه‌گیری ریسک ابتلا به شبیه درجه‌ی مرکزیت را به جریان می‌اندازد که یک اقدام محلی است که تنها بخش کوچکی از کل شبکه را در نظر می‌گیرد. ما چنین گزینه‌هایی را به عنوان کارهای آینده ترک کردیم. اندازه‌گیری مقام آستانه خطی جدید می‌تواند به این عنوان تفسیر شود که چه مقدار عامل i می‌تواند نفوذ خود را در یک شبکه گسترش دهد و چه مقدار سرمایه‌گذاری منابع خارجی سیستم رسمی بدون توجه به جهت‌های لبه قادر به متقاعد کردن همسایگان خود است. از نقطه نظر مثبت، این سرمایه‌گذاری منابع می‌تواند ظرفیت عامل i را برای مدیریت مخاطبین خود در شبکه نشان دهد. از نقطه نظر منفی یا مشکوک، این می‌تواند بیانگر توانایی او در رشوه باشد. به هر حال این اندازه‌گیری وجود

بگیرید. فرض کنید $F_t(X) \subseteq V$ مجموعه‌ای از گره‌ها باشد که در تکرار t فعال می‌شود. در ابتدا در گام $t = 0$ ، تنها گره‌ی X فعال باشد، به این معنا که $F_0(X) = X$ در $t+1$ تکرار بعدی، گره‌ی $i \in V$ فعال می‌شود اگر و تنها اگر شرط زیر برقرار باشد:

$$\sum_{j \in F_t(X)} w_{ji} \geq f(i)$$

این فرآیند زمانی که هیچ فعال سازی دیگری رخ ندهد متوقف می‌شود. به عبارت دیگر، یک گره‌ی i زمانی فعال می‌شود که مجموع وزن‌های گره‌های فعال متصل به آن بزرگ‌تر یا مساوی مقاومت آن در مقابل تأثیرپذیری باشد. **تعریف ۲:** فرض کنید (G, w, f) یک گراف تأثیر باشد، انتشار تأثیر X برابر است با:

$$F(X) = \bigcup_{t=0}^k F_t(X) = f_0(X) \cup \dots \cup F_k(X)$$

که در آن، $k = \min \{t \in \mathbb{N} | F_t(X) = F_{t+1}(X)\} \leq n$ مقدار t در $F_t(X)$ سطح فعلی انتشار را نشان می‌دهد. توجه کنید که برای هر مجموعه‌ی فعال ساز X ، انتشار تأثیر $F(X)$ در زمان چند جمله‌ای قابل محاسبه است.

سنجه‌های جدید مرکزیت و مرکزی سازی

اقدامات جدید مرکزیت و مرکزی سازی در این بخش اقدامات جدید نفوذ بر اساس گسترش نفوذ تحت مدل آستانه‌ی خطی بررسی می‌شود.

آستانه‌ی خطی مرکزیت: مدل آستانه‌ی خطی به جای مدل آشنای مستقل به هیچ‌گونه شانسی بستگی ندارد بلکه به ظرفیت نفوذ هر عامل و مقاومت آن‌ها در برابر نفوذ وابسته است. ما یک اندازه‌گیری مرکزی برای رتبه‌بندی کاربران شبکه به صورت زیر معرفی می‌کنیم.

تعریف ۳: اجازه دهید (G, w, f) گراف نفوذ با $G = (V, E)$ و $|V| = n$ و a عضو V یک عامل باشد. رتبه‌بندی آستانه‌ی خطی a به صورت زیر مشخص می‌شود:

$$LTR(i) = \frac{|F(\{i\}) \cap \text{Neighbors}(i)|}{n}$$

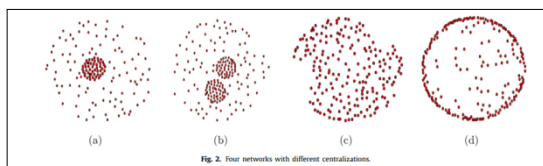
تعریف ۴: اجازه دهید (G, w, f) گراف نفوذ با $G=(V, E)$ و $n=|V|$. مرکزی سازی آستانه ی خطی به صورت زیر تعریف می شود:

$$LTC(G) = \frac{|F(\hat{C}(G))|}{n}$$

where $\hat{C}(G) = \{i \in V | i \text{ belongs to the main core of } G\}$

توجیه این اندازه گیری کاملاً طبیعی است. همانطور که بازیگران خارج از k -shell درجه ای کوچکتر از بازیگران داخلی دارند، اولین افراد بیشتر می توانند با دومین ها تحت تاثیر قرار بگیرند.

در مثالی دیگر، چنین اندازه گیری را به چهار شبکه اعمال شده است. شکل شماره (۱) نتایجی را که تقریباً تجسم تجزیه معمول k -shell را به تصویر می کشد، یعنی گره ها به گونه ای توزیع شده اند که یک گره دارای درجه بزرگتری است که نزدیک تر به مرکز است. توجه داشته باشید که گراف a در مرکز بیشتر مرکزی سازی شده است؛ بنابراین فعال سازی اولیه ی یک مجموعه ی بزرگ است که قادر است نفوذ خود را بر گره های دور منتشر کند. از این رو، اگر گره ها در مرکز توانایی نفوذ خود را بر گره های دور فراهم آورند، این می تواند نمونه خوبی از یک گراف با یک مرکز گرادین خطی بالا باشد. گراف (b) دو هسته بزرگ را به جای یکی نشان می دهد. اگرچه مرکزی سازی بالایی دارد، گره هایی از سطح هستند که فقط از یک هسته از دو هسته قابل دسترسی هستند؛ بنابراین اندازه گیری مای تواند مقادیر پایین تری را نسبت به نمودار اول به دست آورد. نهایتاً به عنوان هسته اصلی هر دو گراف (c) و (d) پایین تر است، فعال سازی اولیه ممکن است برای گسترش نفوذ خود را به گره های محیطی کافی نباشد و در نتیجه مقادیر کمتری برای اندازه گیری مرکزی سازی به دست آید.



شکل (۱): نتایج تجسم تجزیه معمول k -shell

دو شبکه ی مختلف را نشان می دهد. یک شبکه ی رسمی روابط شناخته شده ی بین فردی و یک شبکه ی غیررسمی که عوامل می توانند در شبکه ی رسمی استفاده کنند. این واقعیت بسیار رایجی است که ما دائماً در معرض آن هستیم. به عنوان مثال، در شبکه تویتر، هزاران نفر از کاربران می توانند به روش های مختلفی برای انتشار یک رویداد تعامل داشته باشند. این اخبار بسته به ظرفیت نفوذ هر عامل مورد علاقه در اخبار، از طریق شبکه تحت پوشش قرار می گیرد. اگر چه کاربرانی که دارای بیشترین نفوذ هستند تأثیر بیشتری روی شبکه ایجاد می کنند، هر کاربر آزاد است اگر می خواهد در مورد رویدادی نظر دهد. در این شبکه ی رسمی کاربران نمی توانند به طور مستقیم با تصمیمات سایر کاربران مواجه شوند. با این وجود، سازمان دهندگان این رویداد می توانند از شبکه خارجی خود از مخاطبین (دوستی، رسانه ها و غیره) برای کمک به گسترش این رویداد استفاده کنند؛ بنابراین، مجموعه فعال سازی اولیه معمولاً توسط تنها یک بازیگر تشکیل نمی شود، بلکه توسط چندین مورد که به روشی دیگر که قابل مشاهده در شبکه ی رسمی نیست تشکیل می شود. توجه داشته باشید که افزایش اندازه مقام آستانه خطی جدید به اندازه گره های فعال سازی اولیه بستگی ندارد، همان طور که میزان تأثیر آن در عوامل می تواند باهم به عنوان یک ائتلاف عمل کند. در این اندازه گیری، یک عامل با سهم کمی در گسترش نفوذ ممکن است به دلیل همسایگانش رتبه خوبی داشته باشد. این در گسترش پدیده نفوذ و همچنین در بسیاری از اقدامات مرکزی است که رفتار کل شبکه را در نظر می گیرد. به عنوان مثال، اقدامات مبتنی بر محدوده ی خاصی از منابع مانن رتبه صفحه، جایی که یک عامل تنها می تواند به دلیل ارتباط مستقیم با بازیگران با نفوذ دیگر تأثیر بگذارد.

مرکزی سازی آستانه ی خطی: با استفاده از مدل آستانه ی خطی، ما یک اندازه گیری مرکزی جدید برای تعیین اینکه چه قدر کل شبکه مرکزی سازی شده است را تعریف می کنیم. هسته ی K یک گراف، بیشینه زیرگرافی است که در آن هر راس دارای کمینه درجه ی K است. K -shell زیرگرافی از گره های k -core است، اما نه در $(1+k)$ -core. هسته ی مرکزی، هسته با بیشترین درجه است.

۳- یافته‌های تحقیق

مثالی از یک گراف با آستانه تمرکز خطی بالا. گراف b دو هسته بزرگ را به جای یک هسته نشان می‌دهد. با این وجود، تمرکز زیادی در مرکز دارد، گره (عضو) هایی در سطح وجود دارند که می‌توانند تنها از طریق یک و یا دو هسته در دسترس باشند و بنابراین اندازه گیری ما مقدار کمتری را نسبت به گراف اول بازمی‌گرداند. در نهایت، هنگامی که هسته اصلی هر دو گراف c و d در سطح پایین تر هستند، انرژی فعال سازی اولیه احتمالاً برای تأثیر گذاشتن بر عضو های پیرامونی کافی نخواهند بود که این امر نتیجه اندازه گیری تمرکز در سطوح پایین تر انرژی است.

آزمایشات و نتایج: در این قسمت، نتایج معیار تأثیر مختلف را در چهار شبکه مختلف ارائه شده است:

- شبکه Higgs: شبکه ای بزرگ و جهت دار
- شبکه arXiv: شبکه ای بزرگ و فاقد جهت
- شبکه Dining-table: شبکه کوچک جهت دار
- شبکه Dolphins social network: شبکه کوچک فاقد جهت

ما چهار معیار برای مرکزیت در نظر می‌گیریم: رتبه صفحه، مرکزیت کاتر، رتبه بندی آشاری و مقام آستانه خطی جدید. همچنین LTC و ACC را به عنوان معیارهای متمرکز کردن در نظر گرفته شده است. دو مجموعه شبکه داده اول توسط Collection SNAP's Stanford Large Network Dataset ساخته شده‌اند و مورد سوم برای Pajek datasets در دسترس بوده و مورد آخر نیز برای Data Repository UCI Network. آزمایشات با زبان پایتون ۳ برنامه نویسی شده‌اند. برای کار با گراف‌ها، ما از NetworkX library استفاده شده است.

هر شبکه به عنوان یک گراف تأثیر نشان داده شده است (G, w, f) . برای هر واکنش دهنده $v \in V$ ، ما یک $f(i) = \lfloor \bar{w}/2 \rfloor + 1$ برچسب ایجاد شده است که، بنابراین یک واکنش دهنده تنها زمانی که به $\sum_{(j,i) \in E(G)} w_{ji}$ فعال سازی اولیه تعلق داشته باشد، فعال می‌شود و یا اعضا فعال بیش از نصف کل انرژی تأثیر وارده را به آن وارد می‌کنند.

شبکه Higgs

اولین مجموعه داده که همه ی توپیت های مرتبط با

آزمایش‌ها Higgs boson که بین اول تا هفتم جولای ری-توپیست شده بوده است را شامل می‌شود. این مجموعه نخست برای پروسه توسعه اطلاعات بر روی توپیتر استفاده می‌شده است. این مجموعه داده اجازه می‌دهد تا یک گراف تأثیر با $|V|=n=255491$ واکنش دهنده و $|E|=328132$ یال جهت دار تولید کنیم. یک یال جهت دار $E \ni (i, z)$ نشان دهنده یک واکنش دهنده i است که ری-توپیست می‌کند واکنش دهنده z را. بنابراین می‌توان گفت که i تغییر معینی را بر اعمال می‌کند. وزن یال w_{ij} نشان دهنده آن است که چند بار واکنش دهنده z واکنش دهنده i را ری-توپیست کرده است.

طبق تعریف، معیارهای رتبه بندی آشاری و مقام آستانه خطی جدید هر دو همگرا می‌شوند. برای محاسبه معیار رتبه بندی آشاری، ما از احتمال $p = 0.1$ استفاده می‌کنیم. برای محاسبه رتبه صفحه از عامل محرک $\alpha = 0.85$ که به صورت پیش فرض گذاشته شده استفاده می‌کنیم و الگوریتم همگرا شده قبل از ۱۰۰ تکرار نیز به صورت پیش فرض گذاشته می‌شوند. برای محاسبه مرکزیت کاتب نیز از مقادیر $\alpha = 0.1$ و $\beta = 0.1$ استفاده می‌کنیم که آن‌ها نیز به صورت پیش فرض انتخاب شده‌اند. هر چند که به جای رتبه صفحه، با مقادیر ترانس استاندارد و حتی یک میلیون تکرار، الگوریتم کاتب واگرا شده است؛ بنابراین برای آنکه همگرایی بتواند مقدار مناسب و عملی از منابع زمان و حافظه استفاده کند، ما تصمیم گرفتیم تا ترانس را از استاندارد قابل توجه ۶ رقم به تنها ۲ رقم کاهش دهیم. قبل از مقایسه نتایج معیارهای تأثیر مختلف، اجازه بدهید تا بر روی معیار مقام آستانه خطی جدید متمرکز شویم. جدول شماره (۱) و شکل شماره (۲) نشان می‌دهند که چه تعداد واکنش دهنده به سطوح گسترش برابر ۰، ۱ و ۲ و... رسیده‌اند. توجه کنید که تقریباً ۹۰ درصد واکنش دهنده‌ها به چهارمین سطح گسترش نمی‌رسند. هر چند که منحنی شکل گرفته توسط این دو متغیر کاهش شدید یکنواختی ندارد. در حقیقت، رایج ترین واکنش دهنده‌ها آن‌هایی هستند که دقیقه به دو سطح گسترش رسیده‌اند و هیچ واکنش دهنده‌ای وجود ندارد که دقیقه یازده سطح گسترش را طی کرده باشد. علاوه بر آن، بیشترین سطح گسترش واکنش دهنده‌ها که با معیار مقام آستانه خطی

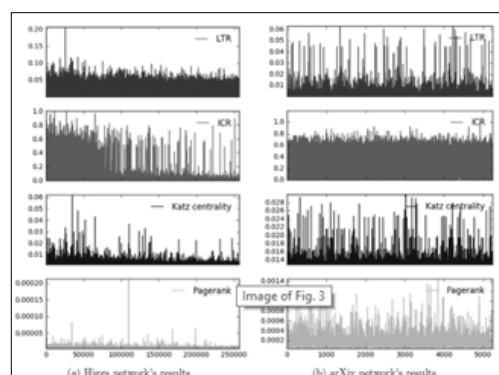
۴- نتیجه‌گیری

این مقاله بر روی مقایسه سنج‌های مرکزیت که مشتق شده از مدل آستانه خطی بود متمرکز شده است. در شبکه‌های اجتماعی مختلف بررسی شد. رتبه‌بندی آستانه خطی را معرفی شده است؛ یک سنج مرکزیت جدید مبتنی بر مدل آستانه خطی که برای هر عضو، به عنوان تعداد اعضایی در نظر گرفته می‌شد که می‌توانست انتشار نفوذ صورت دهد، ارائه شد.

این سنج را با سه سنج مرکزیت دیگر (مرکزیت کاتر، رتبه صفحه و رتبه‌بندی آبشاری مستقل) بر اساس نفوذ و تأثیر در ۴ شبکه واقعی دو شبکه بزرگ (یکی جهت‌دار، دیگری بی‌جهت) و دو شبکه کوچک (یکی جهت‌دار، دیگری بی‌جهت) مقایسه شده است. شبکه بزرگ‌تر ۲۵۵۴۹۱ عضو و ۳۲۸۱۳۲ ارتباط بین اعضا را دارا بود. برای شبکه بزرگ، حتی برای شبکه کوچک جهت‌دار، در ادامه نشان داده شد که همبستگی مابین سنج‌ها، پایین است که به این معنی است که همه آن‌ها حوزه نفوذ مختلفی را فراهم می‌کنند و در نهایت، نتایج گوناگونی را بروز می‌دهند. به طور کلی، یافته‌ها نشان می‌دهد که مقدار حقیقی از نتایج همبستگی در ضریب اسپیرمن بالاتر از ضریب کندال است. علاوه بر این، رتبه‌بندی آستانه خطی با رتبه‌بندی مستقل آبشار، بیشترین میزان انحراف را حاصل می‌دهد. برای مقایسه این سنج، رتبه‌بندی آستانه خطی مقادیر مختلف بیشتری را ایجاد می‌کند. این نشان می‌دهد که رتبه‌بندی آستانه خطی، بسیار پرکاربردتر برای رتبه‌بندی اعضا در یک شبکه خواهد بود. در کنار سنج مرکزیت، در این مقاله یک سنج مرکزی سازی به نام مرکزی سازی آستانه خطی را معرفی شده است؛ که به تعداد اعضایی که می‌توانستند توسط گره‌هایی که به هسته اصلی شبکه تعلق داشتند تحت تأثیر قرار بگیرند تطابق داشت. اثبات شد که این سنج برای شبکه‌های با تعداد زیاد اعضا بسیار مفید واقع خواهد شد. همچنین این سنج با ضریب خوشه‌بندی میانگین نیز مقایسه شد و نتیجه گرفته شد که هرکدام از این سنج‌ها، مرکزی سازی متفاوتی را ایجاد می‌کنند و برای استخراج اطلاعات متفاوتی از شبکه کاربرد دارند.

جدید مرتبط می‌شوند ($\tau = 0.77$ and $\rho = 0.89$). این دین معناست که برای این شبکه، گره‌ها (اعضا) می‌توانند امتیاز مقام آستانه خطی جدید بالایی را با رسیدن به سطح گسترش پایین داشته باشند.

جدول (۱): تعداد واکنش دهنده به سطوح گسترش	
Max t	#actors
۰	۴۲۱۲۹
۱	۶۳۱۸۹
۲	۷۳۲۶۰
۳	۵۱۱۱۸
۴	۱۱۵۱۶
۵	۱۴۸۶۱
۶	۲۱۱
۷	۱۷۱
۸	۲۲
۹	۸
۱۰	۴
۱۲	۲



شکل (۲): تعداد واکنش دهنده به سطوح گسترش

پيشنهادات و كارهاى آينده

براى كارهاى آتى، مى‌توان از مركزى سازى فريمن براى برخى از سنجه‌هاى مركزيت مبتنى بر حوزه انتشار به منظور

مقايسه هر چه بيشتر سنجه‌هاى مركزيت با روش مركزى سازى آستانه خطى استفاده كرد.

منابع

- 5-Adamic, L., & Adar, E. (2005). How to search a social network. *Social networks*, 27(3), 187-203.
- 6-Altaf-Ul-Amine, M., Nishikata, K., Korna, T., Miyasato, T., Shinbo, Y., Arifuzzaman, M., ... & Kanaya, S. (2003). Prediction of protein functions based on k-cores of protein-protein interaction networks and amino acid sequences. *Genome Informatics*, 14, 498-499.
- 7-Bader, G. D., & Hogue, C. W. (2003). An automated method for finding molecular complexes in large protein interaction networks. *BMC bioinformatics*, 4(1), 1-27.
- 8-Bavelas, A. (1948). A mathematical model for group structures. *Human organization*, 7(3), 16-30.
- 9-Bollobás, B. (1984). The evolution of random graphs. *Transactions of the American Mathematical Society*, 286(1), 257-274.
- 10-Bonacich, P. (1972). Factoring and weighting approaches to status scores and clique

- identification. *Journal of mathematical sociology*, 2(1), 113-120.
- 11-Corr, P. J., & Matthews, G. E. (2020). *The Cambridge handbook of personality psychology*. Cambridge University Press.
- 12-Domingos, P., & Richardson, M. (2001, August). Mining the network value of customers. In *Proceedings of the seventh ACM SIGKDD international conference on Knowledge discovery and data mining* (pp. 57-66).
- 13-Easley, D., & Kleinberg, J. (2010). *Networks, crowds, and markets: Reasoning about a highly connected world*. Cambridge university press.
- 14-Foster, K. C., Muth, S. Q., Potterat, J. J., & Rothenberg, R. B. (2001). A faster katz status score algorithm. *Computational & Mathematical Organization Theory*, 7, 275-285.
- 15-Freeman, L. C. (2002). Centrality in social networks: Conceptual clarification. *Social network: critical concepts in sociology. Londres: Routledge*, 1, 238-263.
- 16-Gaertler, M., & Patrignani, M. (2004, March). Dynamic analysis of the autonomous system graph. In *IPS 2004, International Workshop on Inter-domain Performance and Simulation, Budapest, Hungary* (pp. 13-24).
- 17-Goldenberg, J., Libai, B., & Muller, E. (2001). Using complex systems analysis to advance marketing theory development: Modeling heterogeneity effects on new product growth through stochastic cellular automata. *Academy of Marketing Science Review*, 9(3), 1-18.
- 18-Gómez, D., Figueira, J. R., & Eusébio, A. (2013). Modeling centrality measures in social network analysis using bi-criteria network flow optimization problems. *European Journal of Operational Research*, 226(2), 354-365.
- 19-Granovetter, M. (1978). Threshold models of collective behavior. *American journal of sociology*, 83(6), 1420-1443.
- 20-Gruhl, D., Guha, R., Liben-Nowell, D., & Tomkins, A. (2004, May). Information diffusion through blogspace. In *Proceedings of the 13th international conference on World Wide Web* (pp. 491-501).
- 21-Katz, L. (1953). A new status index derived from sociometric analysis. *Psychometrika*, 18(1), 39-43.
- 22-Kempe, D., Kleinberg, J. M., & Tardos, É. (2005, July). Influential nodes in a diffusion model for social networks. In *ICALP* (Vol. 5, pp. 1127-1138).
- 23-Kendall, M. G. (1938). A new measure of rank correlation. *Biometrika*, 30(1/2), 81-93.
- 24-Kitsak, M., Gallos, L. K., Havlin, S., Liljeros, F., Muchnik, L., Stanley, H. E., & Makse, H. A. (2010). Identification of influential spreaders in complex networks. *Nature physics*, 6(11), 888-893.
- 25-Langville, A. N., & Meyer, C. D. (2004). Deeper inside pagerank. *Internet Mathematics*, 1(3), 335-380.
- 26-Li, J., Peng, W., Li, T., Sun, T., Li, Q., & Xu, J. (2014). Social network user influence sense-making and dynamics prediction. *Expert Systems with Applications*, 41(11), 5115-5124.
- 27-Liu, W., Yue, K., Wu, H., Li, J., Liu, D., & Tang, D. (2016). Containment of competitive influence spread in social networks. *Knowledge-Based Systems*, 109, 266-275.
- 28-Lozares, C., López-Roldán, P., Bolibar, M., & Muntanyola, D. (2015). The structure of global centrality measures. *International Journal of Social Research Methodology*, 18(2), 209-226.
- 29-Morone, F., & Makse, H. A. (2015). Influence maximization in complex networks through optimal percolation. *Nature*, 524(7563), 65-68.
- 30-Page, L., Brin, S., Motwani, R., & Winograd, T. (1999). *The PageRank citation ranking: Bringing order to the web*. Stanford InfoLab.
- 31-Riquelme, F., Gonzalez-Cantergiani, P., Molinero, X., & Serna, M. (2018). Centrality measure in social networks based on linear threshold model. *Knowledge-Based Systems*, 140, 92-102.
- 32-Saramäki, J., Kivelä, M., Onnela, J. P., Kaski, K., & Kertesz, J. (2007). Generalizations of the clustering coefficient to weighted complex networks. *Physical Review E*, 75(2), 027105.
- 33-Schelling, T. C. (2006). *Micromotives and macrobehavior*. WW Norton & Company.
- 34-Seidman, S. B. (1983). Network structure and minimum degree. *Social networks*, 5(3), 269-287.
- 35-Song, X., Tseng, B. L., Lin, C. Y., & Sun, M. T. (2006, August). Personalized recommendation driven

- by information flow. In *Proceedings of the 29th annual international ACM SIGIR conference on Research and development in information retrieval* (pp. 509-516).
- 36-Sun, J., & Tang, J. (2011). A survey of models and algorithms for social influence analysis. *Social network data analytics*, 177-214.
- 37-Wang, L., Ma, L., Wang, C., Xie, N. G., Koh, J. M., & Cheong, K. H. (2021). Identifying influential spreaders in social networks through discrete moth-flame optimization. *IEEE Transactions on Evolutionary Computation*, 25(6), 1091-1102.
- 38-Watts, D. J., & Strogatz, S. H. (1998). Collective dynamics of 'small-world' networks. *nature*, 393(6684), 440-442.
- 39-Ye, S., & Wu, F. (2013). Measuring message propagation and social influence on Twitter. com. *International Journal of Communication Networks and Distributed Systems*, 11(1), 59-76.
- 40-Zhang, X., Zhu, J., Wang, Q., & Zhao, H. (2013). Identifying influential nodes in complex networks with community structure. *Knowledge-Based Systems*, 42, 74-84.

©Authors, Published by Journal of Intelligent Knowledge Exploration and Processing. This is an open-access paper distributed under the CC BY (license <http://creativecommons.org/licenses/by/4.0/>).



مقاله پژوهشی

نوگرایی مصرف کننده و عملکرد برندهای آنلاین با توجه به نقش میانجی تعامل مشتریان با برند آنلاین

(مورد مطالعه فروشگاه اوکالا در شهر مشهد)

Doi: 10.30508/KDIP.2022.366246.1052

علی میرزایی راد^۱ | محمدحسین همایونی راد^۲ | سعیده باباجانی محمدی^۳

۱- کارشناس ارشد مدیریت بازرگانی، موسسه آموزش عالی فردوس، مشهد، ایران.

۲- استادیار گروه مدیریت، موسسه آموزش عالی سناباد، گلپه‌ار، ایران.

۳- استادیار گروه مدیریت، موسسه آموزش عالی فردوس، مشهد، ایران.

تاریخ دریافت: ۱۴۰۱/۰۷/۲۷

تاریخ پذیرش: ۱۴۰۱/۰۹/۱۷

صفحه: ۵۵ - ۵۰

چکیده:

امروز دیگر بسیاری از محققان و شرکت‌های فعال در حوزه تجارت برخط به این نتیجه رسیده‌اند که با ارزش‌ترین دارایی یک شرکت در جهت بهبود فرایند بازاریابی، برند و عملکرد برند است. با این وجود، پذیرش برندها در محیط آنلاین یکی از مهم‌ترین مباحث در حوزه سامانه‌های اطلاعاتی و برند بوده که کمتر مورد توجه قرار گرفته است. از این رو محقق در این تحقیق به بررسی رابطه عوامل مؤثر بر نوگرایی مصرف کننده و عملکرد برندهای آنلاین با توجه به نقش میانجی تعامل مشتریان با برند آنلاین پرداخته است. با توجه به هدف و فرضیه‌ها، تحقیق حاضر از نوع استفاده کاربردی و از حیث روش تحلیل و آزمون فرضیه‌ها توصیفی و همبستگی است. جامعه آماری شامل کلیه شهروندان مشهدی هستند که از اوکالا به صورت آنلاین خرید می‌کنند با توجه به ویژگی‌های جامعه آماری و شیوه نمونه‌گیری تصادفی ساده در دسترس و استفاده از روش مورگان، ۳۸۴ واحد نمونه در نظر گرفته شد. برای جمع‌آوری اطلاعات از پرسشنامه استاندارد برایان و همکاران (۲۰۲۰) استفاده و با استفاده از روایی صوری و استفاده از ضریب اعتماد آلفای کرونباخ، ضریب پایایی بالای ۰/۸ برای آن محاسبه گردید. آزمون معادلات ساختاری برای پاسخ به فرضیه‌های تحقیق استفاده شد. نتایج نشان داد، تعامل مشتری با برند، رابطه بین عوامل مؤثر بر نوگرایی مصرف کننده با عملکرد برند در فضای آنلاین را میانجی‌گری می‌کند. بین عوامل مؤثر بر نوگرایی مصرف کننده و عملکرد برند در فضای آنلاین، و همچنین میان عوامل مؤثر بر نوگرایی مصرف کننده و میزان تعامل مشتری با برند رابطه مثبت و معنی‌دار وجود دارد. بین میزان تعامل مشتری با برند و عملکرد برند در فضای آنلاین رابطه مثبت و معنی‌دار وجود دارد. لذا برای ایجاد روابط مؤثر و بلندمدت، شناخت لایه‌های زیرین اثرگذار در روابط مصرف کننده ضروری است و درک ناکافی از پیوندهای احساسی تمامی هزینه‌های بازاریابی صرف شده شرکت را از بین می‌برد.

کلمات کلیدی: نوگرایی مصرف کننده، عملکرد برند، برندهای آنلاین، تعامل مشتریان، تعامل با برند آنلاین.

۱- مقدمه

خرید آنلاین در دنیا نسبت به قبل مقرون به صرفه‌تر و راحت‌تر می‌شود و به مصرف کنندگان اجازه می‌دهد قیمت‌ها، کیفیت و خدمات کالاها را به روشی مطلوب مقایسه کنند. به این ترتیب رضایت مشتریان از بین می‌رود و اغلب ترجیح می‌دهند خرید خود را به صورت آنلاین انجام دهند. فروشگاه‌ها و تولیدکنندگان بزرگ و پیشرو در بازاریابی و تجارت آنلاین معمولاً برای موفقیت باید ابتدا ارزش و ریسک درک شده را برای مشتری ایجاد کنند. مشتریان در فضای آفلاین و آنلاین به راحتی هر فروشگاه را با فروشگاه‌های دیگر مقایسه می‌کنند تا بتوانند محصول یا خدمات مورد نیاز خود را با بهترین کیفیت و کمترین قیمت دریافت کنند و سپس اقدام به خرید کنند (همیلتون^۱، ۲۰۱۹). امروزه آنچه برای مشتریان مهم است استفاده از محصولات و خدماتی است که باعث صرفه‌جویی در زمان و افزایش راحتی مشتری می‌شود. تمایل به دریافت یا دریافت خدمات از سازمان‌هایی که به آسایش مشتری اهمیت می‌دهند، ناشی از تمایل مشتری به سبک زندگی راحت‌تر است و سهولت استفاده از خدمات، شاخصی کلیدی برای آنها محسوب می‌شود. درک مشتری از سهولت استفاده از خدمات بر سنجش مشتری از خدمات و رضایت او، اعتبار خدمات درک شده، اعتماد و وفاداری او به برند تأثیر می‌گذارد. اهمیت هر یک از ویژگی‌های سهولت درک و استفاده از خدمات از نظر کیفیت مشتری بسیار متغیر است، مشتریانی که به صرفه‌جویی در زمان اهمیت می‌دهند اعتبار بیشتری برای سهولت استفاده از خدمات قائل هستند. مأموریت اصلی ارائه دهندگان خدمات افزایش خروجی مثبتی است که مشتری دریافت می‌کند. مهم‌ترین یافته‌ها رضایت مثبت مشتری و رابطه مثبت بین رضایت مشتری و وفاداری است (احمدی^۲، ۲۰۱۹).

همچنین تعامل مشتریان با شرکت‌ها از حالت سنتی و یک طرفه خارج شده و فضای مجازی و شبکه‌های اجتماعی اهمیت بیشتری پیدا کرده است. امروزه شرکت‌ها از فناوری‌های مجازی برای فعالیت‌های خود و همچنین برای

بهبود کارایی، بهره‌وری، هزینه یا کیفیت عملیات مشتریان خود استفاده می‌کنند (آلن، جوجیمس، و وو^۳، ۲۰۲۲). در حال حاضر مشتریان انتظار دارند برندها یا کالاهایی که می‌خرند، آنها را راضی کند؛ اما رضایت شرط کافی برای ایجاد یک رابطه مستمر با برند نیست. مشتریان در پی ایجاد پیوندهایی فراتر از رضایت‌اند. پیوندهایی که مبتنی بر وابستگی عاطفی است. در سال‌های اخیر، اشتیاق به برند یکی از موضوع‌های مهم در حوزه مطالعات برند بوده است. به عبارتی، زمانی که یک برند به شکلی عمل می‌کند که توانایی ارضای نیازهای واقعی و ملموس از یک طرف و خواسته‌های احساسی و غیرملموس از طرف دیگر را داشته باشد، می‌تواند احساسی قوی در مصرف‌کننده ایجاد کند که می‌توان آن را به اشتیاق تعبیر کرد و این گونه و رفته‌رفته مفهوم اشتیاق برند در تحقیق‌های مصرف‌کننده، به طور گسترده‌ای در قیاس با اشتیاق انسانی مفهوم‌سازی گردیده است. اما تحولات جدید در حوزه فناوری اطلاعات، فضاهای مجازی جذابی مانند شبکه‌های اجتماعی را ایجاد کرده است که روز به روز در حال گسترش هستند و عرصه را برای تبلیغات کالاها و خدمات تولیدکنندگان فراهم می‌کنند. با پیشرفت فناوری اینترنت، شرکت‌ها از شبکه‌های اجتماعی برای تبلیغ و انتشار اطلاعات برند خود استفاده می‌کنند. رسانه‌های اجتماعی ارتباطات سنتی بازاریابی را تغییر داده‌اند. کاربران اینترنت به تدریج در حال شکل‌گیری روابط تجاری هستند که به طور سنتی توسط بازاریابان شکل گرفته است. در اکثر کشورهای پیشرفته دنیا شبکه‌های اجتماعی بسیار مورد استفاده قرار گرفته‌اند و تقریباً تمام ابعاد زندگی مردم را پوشش داده‌اند. سازمان‌دهندگان این شبکه‌ها، توانسته‌اند بهترین استفاده را از این ابزارها ببرند و تولیدکنندگان محصولات و خدمات توانستند بر اساس اطلاعاتی که از این شبکه‌ها دریافت می‌کنند، اعتماد و وفاداری به برند را در مشتری تقویت نمایند (میچل، و بالابانیس^۴، ۲۰۲۱؛ ویکرس^۵، ۲۰۱۷). مقوله عملکرد برند در دنیای امروزی طرفداران زیادی در دنیای بازاریابی پیدا کرده است. در بخش خدمات، برندها

1- Hamilton

2- Ahmadi

3- Allen, Jochims, & Wu

4- Mitchell, & Balabanis

5- Vickers

راهی سریع برای شناسایی و متمایز کردن خود و ایجاد تصویری در ذهن مشتریان هستند. عملکرد برند مهم‌ترین چیزی است که بیشتر استراتژی‌های بازاریابی روی آن تمرکز می‌کنند و تمایل دارند آن را برجسته کنند. برای این منظور بخش‌های خدماتی سعی می‌کنند با تأثیری که بر ادراک مشتریان از خدمات دریافتی می‌گذارند، با مشتریان ارتباط برقرار کرده و تصویری مطلوب در ذهن مشتری ایجاد کنند (جیانگ، جانگ، جوو و کیم، ۲۰۲۱). امروزه سازمان‌ها به دنبال جذب و حفظ مشتریان بیشتر هستند. عملکردی که برند در ذهن مشتریان ایفا می‌کند یکی از ساختارهای برجسته در شکل‌گیری وفاداری به برند و عامل موفقیت یک سازمان است. سازمان‌ها و شرکت‌ها می‌توانند با افزایش مطلوبیت خدمات، محصولات، امکانات، فضای فروشگاه و غیره، تصویر ذهنی مثبت مشتریان از محصول یا برند خود را بهبود بخشند و از آن به عنوان یک مزیت رقابتی عمده استفاده کنند (واندرسچی، پنتلیرو و دال، ۲۰۲۰). اکنون بسیاری از تولیدکنندگان بر این باورند که تمرکز بر برند و عملکرد برند مهم‌تر از کار بر روی خطوط تولید است. زیرا هر چه قدر هم که در تولید قدرتمند باشند، اگر نتوانند برند برتری داشته باشند، محصولات و خدماتشان جایی در بازار پیدا نمی‌کند. برندسازی فعالیتی است که در آن تصویری از برند در قلب و ذهن مشتری ایجاد می‌شود. در فرآیند برندسازی سعی می‌شود با استفاده از یک نام، نماد یا لوگو بر ادراک افراد تأثیر گذاشته و هویت برند را در خاطرات تثبیت کند. متأسفانه کاری که توسط شرکت‌های این حوزه در ایران انجام می‌شود تنها بر اساس تصور برند به عنوان یک برند است و تنها راه حل آن تبلیغات است. با این حال، هر چیزی که پیامی را به بازار برساند می‌تواند در این بخش مفید و مثمر ثمر باشد. با توجه به اینکه بیشتر تجارت الکترونیک به صورت آنلاین انجام می‌شود، بازاریابی دیجیتال در این مشاغل بسیار مهم است (سرگزی، ۱۳۹۸). در مبنای نظری رفتار مصرف‌کننده، موردی که همیشه بیشترین توجه را به خود جلب می‌کند، ارزش دریافتی مشتری از محصول یا خدمات است. ارزش واقعی تضاد بین

آنچه مشتری دریافت می‌کند (یعنی کیفیت خدمات، مزایا) و آنچه که برای به دست آوردن آن مزایا از دست می‌دهد. به عبارت دیگر، ارزش ادراکی توسط مشتری یک قضاوت کلی در مورد منافع به دست آمده و هزینه‌ها است. ارزش دریافتی مشتری نقش بسیار مهمی در تعیین رفتار مشتری در انتخاب و قصد انتخاب او دارد و ترجیح مصرف‌کننده برای خرید را شکل می‌دهد (جین و کانگ، ۲۰۱۱). درک رفتار خرید مصرف‌کننده جنبه حیاتی بازاریابی است. رفتار خرید مصرف‌کننده فرآیند تصمیم‌گیری در مورد اینکه چه چیزی بخرد، چه چیزی بخواهد، چه چیزی مصرف‌کننده نیاز دارد و از سوی دیگر عملکرد یک محصول، خدمات یا شرکت است. به درک رفتار مصرف‌کننده کمک می‌کند، بدانیم مشتریان بالقوه چگونه به یک محصول یا خدمات جدید پاسخ خواهند داد و به شرکت یا سازمان کمک می‌کند تا درخواست‌هایی را که توسط مصرف‌کنندگان برآورده نشده است شناسایی کند (لیدیوا و همکاران، ۲۰۱۸).

بر اساس تحقیقات پانتانو، دنیس، و دی پیترو^۵ (۲۰۲۱)، عوامل مؤثر بر کاهش عملکرد برند در فضای آنلاین عبارتند از: تجربه منفی گذشته و ناسازگاری نمادین و ناسازگاری ایدئولوژیک. بنابراین واضح است که عملکرد بد و ضعیف برند تصویری ناخواسته و نامطلوب در ذهن مصرف‌کنندگان به جا می‌گذارد و اینجاست که مفهوم ناسازگاری نمادین مطرح می‌شود. امروزه بسیاری از محققین و شرکت‌ها در زمینه تجارت آنلاین به این نتیجه رسیده‌اند که با ارزش‌ترین دارایی یک شرکت به منظور بهبود فرآیند بازاریابی، برند و عملکرد برند است. با این وجود، پذیرش برندها در محیط آنلاین یکی از مهم‌ترین موضوعات در حوزه سیستم‌های اطلاعاتی و برندها است که کمتر مورد توجه قرار گرفته است. علیرغم مزایای نسبی فراوان تجارت آنلاین، در ایران و برخی کشورهای در حال توسعه چندان مورد استقبال قرار نگرفته است. متأسفانه به دلیل ساختار ضعیف اینترنت در ایران، مردم کمتر از سایر کشورهای پیشرفته از آن استفاده می‌کنند. یکی از دلایل عدم خرید آنلاین، عدم پذیرش ریسک است. تحقیقات در سراسر جهان در مورد موانع

1- Jiang, Yang, Joo, & Kim

2- Vander Schee, Peltier, & Dahl

3- Jin & Kang

4- Lebedeva, et al

5- Pantano, Dennis, & De Pietro

لوکس برای بازاريابان کالاهای لوکس مفید است. امروز دیگر بسیاری از محققان و شرکت‌های فعال در حوزه تجارت برخط به این نتیجه رسیده‌اند که باارزش‌ترین دارایی یک شرکت در جهت بهبود فرایند بازاریابی، برند و عملکرد برند است. با این وجود، پذیرش برندها در محیط آنلاین یکی از مهم‌ترین مباحث در حوزه سامانه‌های اطلاعاتی و برند بوده که کمتر مورد توجه قرار گرفته است. با وجود مزایای نسبی بسیار زیاد تجارت آنلاین و برخط، در ایران و برخی کشورهای درحال توسعه استقبال زیادی از آن نشده است. متأسفانه به دلیل ضعف در ساختار اینترنت در ایران افراد از این مهم نسبت به سایر کشورهای پیشرفته کمتر استفاده می‌کنند. از دلایل کم بودن خرید برخط عدم قبول ریسک است. تحقیق‌ها در سراسر دنیا درباره موانع توسعه تجارت الکترونیکی به اعتماد اشاره کرده‌اند با عنایت به مسائل مطرح شده، به منظور توفیق هر چه بیشتر محیط‌های تجارت الکترونیک اوکالا، محقق به دنبال پاسخ این سوال است که آیا بین نوگرایی مصرف‌کننده با عملکرد برندهای آنلاین با توجه به نقش میانجیگری تعامل مشتریان با برند آنلاین، رابطه معناداری وجود دارد یا خیر؟

۲- مبانی نظری

رفتار مصرف‌کننده: فعالیت‌های افرادی که مستقیماً در تهیه و استفاده از کالاها و خدمات اقتصادی مشارکت دارند، از جمله فرآیندهای تصمیم‌گیری که مقدم بر این فعالیت‌ها و تعیین‌کننده آن است. فرآیند تصمیم‌گیری و فعالیت بدنی که توسط آن افراد در ارزیابی، کسب، استفاده یا رد کالاها و خدمات شرکت می‌کنند (چن، ژو، سوماری و اکسو، ۲۰۱۵).

نوگرایی مصرف‌کننده: عصر حاضر دوران دگرگونی و رقابت بی‌سابقه و تنگاتنگ بین شرکت‌های اقتصادی بین المللی است. شرایط به دلیل از بین رفتن مرزها در بازاریابی، پراکندگی بازارها، کوتاه شدن عمر محصول، تغییرات سریع و روش‌های استفاده از خدمات مشتری و افزایش آگاهی مصرف‌کننده، دائماً در حال تغییر است (یانگ، لین، فیلیر، لیو و ژنگ، ۲۰۲۰). می‌توان ادعا کرد که هر چه دامنه مسائل

توسعه تجارت الکترونیک اعتماد را برجسته کرده است. نوگرایی مصرف‌کننده به عنوان اعتقاد مصرف‌کننده، طوری تعریف شده است که او با تجارت آنلاین از یک سایت خاص سود بیشتری کسب می‌کند. نوگرایی مصرف‌کننده شامل مزایای ارتباط آسان در خرید و همچنین کاهش هزینه‌های معامله است. مصرف‌کنندگان اینترنت آنلاین مزایای خرید آنلاین را بیشتر از خرید سنتی می‌دانند. زیرا نوگرایی مصرف‌کننده، انگیزه اصلی برای خرید آنلاین را در مصرف‌کنندگان فراهم می‌کند (پانتانو، ۲۰۲۱). اگر مصرف‌کننده، از طریق یک وب سایت خاص، ارزش بیشتری در معاملات آنلاین درک کند، معاملات آنلاین و خرید‌ها افزایش می‌یابد. ریسک درک شده تمایل مصرف‌کنندگان را برای خرید کالاها و خدمات آنلاین کاهش می‌دهد (منصوری مویید، مرادی و ملایی، ۱۳۹۷). مطالعات نشان می‌دهد که تقویت نوآوری خرید توسط مشتری، منجر به تمایل خرید دوباره خواهد شد. نوگرایی مصرف‌کننده نیز رابطه مثبت با خرید و تصمیم مصرف‌کننده دارد. در بازار سنتی، قیمت بالا نشانه‌ای از کیفیت بالا است و خواص فیزیکی محصول کاملاً روشن است و مصرف‌کننده می‌تواند آن را ببیند. اما در تجارت الکترونیک، دسترسی آسان به اطلاعات در مورد محصولات جدید مهم است. این اطلاعات، توانایی قضاوت مشتری در خصوص کیفیت محصول را افزایش می‌دهد (سوری، ۱۳۹۹).

مدیران بازاریابی علاقه مند به دانستن قصد مشتریان برای خرید برای افزایش فروش محصولات یا خدمات فعلی خود هستند. بنابراین، اطلاعات مربوط به قصد خرید می‌تواند به مدیران تصمیمات بازاریابی کمک کند که مربوط به تقاضای محصول (محصولات فعلی و جدید)، تقسیم بندی بازار و راهبردهای ارتقاء و ارتقاء است. کالاهای لوکس و انگیزه برای خرید مارک‌های لوکس، برای مصرف‌کنندگان آسیایی به میزان قابل توجهی افزایش یافته است. انگیزه خرید لوکس به طور کامل بر اساس تفکر و بازارهای غربی است. ارزش‌های فرهنگی برای تأثیرگذاری بر رفتار مشتری در بسیاری از مطالعات نشان داده شده است. ایجاد روابط بین ارزش‌های فرهنگی و انگیزه برای مصرف کالاهای

مختلف مانند فشار خون، وزن و دمای بدن ارزیابی می‌کند، یک متخصص بازاریابی با اطلاعاتی در مورد ویژگی‌ها و ابعاد محصول می‌تواند به راحتی سیاست‌های بازاریابی مناسب را ارزیابی کند. تصمیم بگیرید و آنها را اعمال کنید. اما با توجه به پیشینه تحقیق هرگز دیدگاه جامعی وجود ندارد و شاید بتوان گفت معیاری برای سنجش آن وجود ندارد و محققان مختلف معیارهای مختلفی را برای سنجش آن معرفی و استفاده کرده‌اند (بالودس، وورهی، وکلانتو^۵، ۲۰۱۵).

تعامل با برند آنلاین^۲: فروشنده عمدتاً با استفاده از خدمات آنلاین و با استفاده از شبکه‌های مجازی بر اساس برخورد مشتری در فضای مجازی و اینترنتی سعی در جذب مشتری دارد و سعی می‌کند محصول و خدمات خود را با کمترین قیمت و بالاترین امکانات معرفی کند (برایان و همکاران، ۲۰۲۰).

کاردانی ملکی نژاد، خوراکیان، و رحیم نیا (۱۴۰۰)، در تحقیقی با دیدگاه سلسله مراتبی نوجویی مصرف‌کننده به بررسی اثر نوجویی ذاتی بر رفتار نوآورانه مصرف‌کننده با نقش واسط نوجویی به دامنه خاص نتیجه گرفت دیدگاه سلسله مراتبی نوجویی مصرف‌کننده را مورد تأیید است، بویژه اینکه نوجویی شناختی و نوجویی به محصول دارای نوگرایی در زمینه خاص و همچنین نوجویی احساسی و نوجویی به اطلاعات دارای نوگرایی در زمینه خاص بهترین ترکیب‌ها برای پیش‌بینی رفتار نوآورانه مصرف‌کننده می‌باشند. رضایی و حیدرزاده (۱۳۹۹)، در تحقیقی به بررسی نقش تعدیل‌کننده تنوع طلبی، نوآوری و مستعد بودن مصرف‌کننده بر رابطه بین رضایت با وفاداری رفتاری مشتری نشان داد که ضرایب سازه مستقل بر وابسته و تأثیر سازه میانجی بر وابسته به صورت مثبت و معنادار پشتیبانی شده است. نتایج مدل دوم، از ضرایب تأثیر تعاملی سه متغیر تعدیل‌گر با متغیر میانجی بر متغیر وابسته بررسی و نتایج نشان داده است که ضرایب تعاملی دو مسیر به صورت معکوس و معنادار و یک مسیر به صورت مستقیم و معنی‌دار است. بنابراین هر پنج فرضیه

پیرامون رفتار مصرف‌کننده بیشتر باشد، اهمیت آن برای تحقیق بیشتر می‌شود. موضوع خرید یا دریافت خدمات مورد اذعان بسیاری از محققین است. نظریه تصمیم‌گیری انگیزشی برای خرید مصرف‌کننده از منظر مبنای نظری در حوزه تصمیم‌گیری خرید مصرف‌کننده است (نوابی، ۱۳۹۶). این دیدگاه استدلال می‌کند که تصمیم انگیزشی مصرف‌کنندگان برای خرید یا دریافت خدمات به احساسات آنها بستگی دارد. عواطف متعددی مانند لذت، شادی، امید، ترس، تمایلات جنسی و تمایلات شخصی در این زمینه تأیید شده است (منصوری مویید و همکاران، ۱۳۹۷).

عملکرد برند^۱: عملکرد برند معیاری برای سنجش موفقیت برند در بازار است (برایان و همکاران، ۲۰۲۰). برند آنلاین^۲: برند آنلاین یک فن در مدیریت برند تعریف می‌شود که از اینترنت به عنوان واسطه‌ای برای قرار دادن نام تجاری خود در بازار استفاده می‌کند (یزدانی و راکی زاده، ۱۳۹۶). بازاریابی اینترنتی^۳: معامله‌ای است که در آن خرید یا فروش کالا یا خدمات از طریق اینترنت صورت می‌گیرد و منجر به واردات یا صادرات کالا یا خدمات می‌شود (بهانسوری و کریشنا^۴، ۲۰۱۵).

عملکرد برند در فضای آنلاین: موفقیت یک تجارت بدون شک به عملکرد برند آن کسب و کار بستگی دارد. نیاز به سنجش عملکرد سازمان از جنبه‌های مختلف و با توجه به سطوح مختلف اغلب در ادبیات بازاریابی و به عنوان یک متغیر وابسته مورد توجه قرار گرفته است، بنابراین دیدگاه ارزیابی عملکرد از طریق محصولات و خدمات ارائه شده توسط سازمان وجود دارد. به عبارت دیگر، اغلب در بحث برندها، دو سوال اصلی در ذهن ایجاد می‌شود: «چه عواملی باعث استحکام برند می‌شود؟» و «چگونه یک برند قوی بسازیم؟» برای پاسخ به این سوالات، مفهوم گسترده عملکرد برند معرفی شده است. بنابراین مدیران با آگاهی از ابعاد و ویژگی‌های عملکرد برند، مجهزتر شده و می‌توانند استراتژی‌های موثرتری را برای برند اعمال کنند. مانند پزشکی که سلامت بیمار خود را با اندازه‌گیری پارامترهای

- 1- Brand Outcomes
- 2- Online Brand
- 3- Intrnet Marketing
- 4- Bhuvanewari, & Krishnan
- 5- Baldus, Voorhees, & Calantone
- 6- Consumer Engagement

سرمایه‌گذاری افزایش می‌دهد. طراحی / روش / رویکرد - نویسندگان مقالاتی را بررسی می‌کنند که ساختارها و ابعاد فرعی را بررسی می‌کنند که چگونه عوامل مصرف‌کننده بر تعامل برند و نتایج برند تأثیر می‌گذارند. در نهایت، شش نتیجه برند مورد بررسی قرار گرفت: وضعیت برند، تمایل، نگرش، پیوند تأیید و انزجار. مفاهیم عملی - این بررسی از طریق درک عمیق‌تر عواملی که باعث تعامل مصرف‌کننده و عوامل برندی که باعث تعامل مشتری می‌شوند، به ادبیات کمک می‌کند. نویسندگان چارچوبی را ارائه می‌کنند که هم نیازهای تحقیقاتی آینده را شناسایی می‌کند و هم بینش‌هایی را درباره نحوه ایجاد، رشد شرکت‌ها در رابطه با برند مصرف‌کننده ارائه می‌کند.

۳- روش تحقیق

هدف از این پژوهش، بررسی رابطه عوامل موثر بر نوگرایی مصرف‌کننده و عملکرد برندهای آنلاین با توجه به نقش میانجی تعامل مشتریان با برند آنلاین (مورد مطالعه فروشگاه اوکالا در شهر مشهد) بوده است. در جامعه آماری شامل کلیه شهروندان مشهدی هستند که از اوکالا به صورت آنلاین خرید می‌کنند با توجه به ویژگی‌های جامعه آماری و شیوه نمونه‌گیری تصادفی ساده در دسترس و استفاده از روش مورگان، ۳۸۴ واحد نمونه در نظر گرفته شد. محقق برای جمع‌آوری اطلاعات از پرسشنامه ۲۵ سوالی استاندارد برایان و همکاران (۲۰۲۰) استفاده شد. جدول شماره (۱)، منبع سوالات را اشاره نموده است.

تحقیق مورد تأیید قرار گرفته است. دی پلسمارک و همکاران (۲۰۱۸)، در مطالعه‌ای در مورد تأثیر تعامل جامعه با برند مصرف‌کننده آنلاین در زمان بحران آسیب محصول بر پاسخ‌های شناختی مصرف‌کننده و پاسخ‌های رفتاری به تلاش برای بازیابی مؤثرترین نام تجاری. داده‌های اعضای انجمن، برند آنلاین سامسونگ در چین، در طول بحران باتری گلکسی نوت ۷ این برند جمع‌آوری شد. نتایج نشان می‌دهد که تعامل جامعه با برند مصرف‌کننده آنلاین اثر مستقیم و غیرمستقیم بر قصد خرید مجدد از طریق بخش مصرف‌کننده دارد. و اثر آن عمدتاً از طریق مصرف‌کننده غیر مستقیم است. یافته‌ها نشان می‌دهد که سطح بالایی از تعامل با برند توسط مصرف‌کننده می‌تواند پیامدهای منفی ارزش برند را جبران کند. برایان و همکاران (۲۰۲۰) در مطالعه‌ای عوامل مصرف‌کننده را با میزان تعامل و عملکرد برند در فضای آنلاین بررسی کرد. نتایج نشان داد که مصرف‌کنندگان با رسانه‌های اجتماعی راحت هستند و شرکت‌هایی را می‌پذیرند که فضای دیجیتال مشابهی را اشغال می‌کنند. با این حال، برخی از مصرف‌کنندگان ارتباط با شرکت‌های آنلاین را آسان ترمی دانند. پیشینیان چندین عامل مصرف‌کننده هستند و هنوز به طور کامل کشف نشده‌اند. تعامل آنلاین مصرف‌کننده نیز به روش‌های مختلفی تعریف و اندازه‌گیری می‌شود. نتایج برند همچنین بر تعامل مشتری در آینده اصرار دارد. به طور خلاصه، یافته‌های تحقیق در مورد عامل مصرف‌کننده و خطوط تحقیقاتی آینده در مورد نتایج برند، درک تعامل مشتری و استراتژی‌های برندسازی را برای به حداکثر رساندن بازاریابی

جدول (۱): پرسشنامه‌های مورد استفاده در تحقیق			
ردیف	متغیرها	منبع	سوالات مرتبط
۱	نوگرایی مصرف‌کننده	هولوک و همکاران (۲۰۱۴)	۱ الی ۴
۲	عملکرد برند	شیونکی و همکاران (۲۰۱۶)	۵ الی ۱۶
۳	تعامل با برند آنلاین	ویوک و همکاران (۲۰۲۰)	۱۷ الی ۲۵

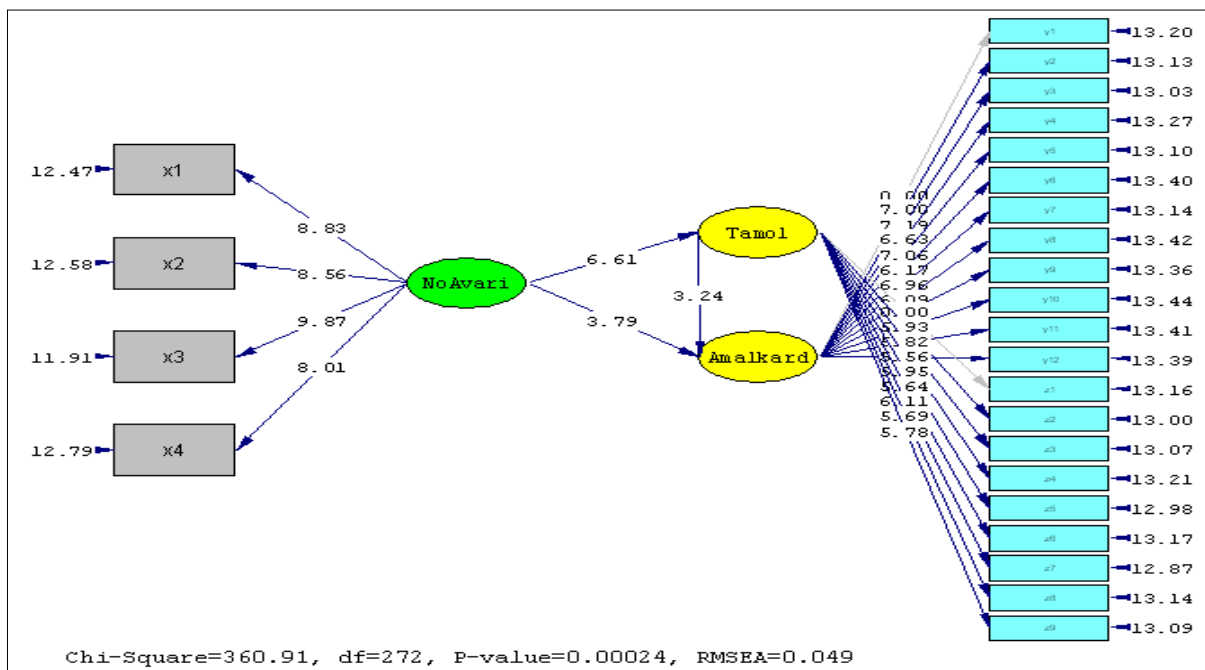
۴- یافته‌های تحقیق

و تحلیل سوالات بخش جمعیت‌شناختی، ویژگی‌های همان گونه که در جدول شماره (۲) بیان شده است، تجزیه جامعه مورد مطالعه را توضیح داده است.

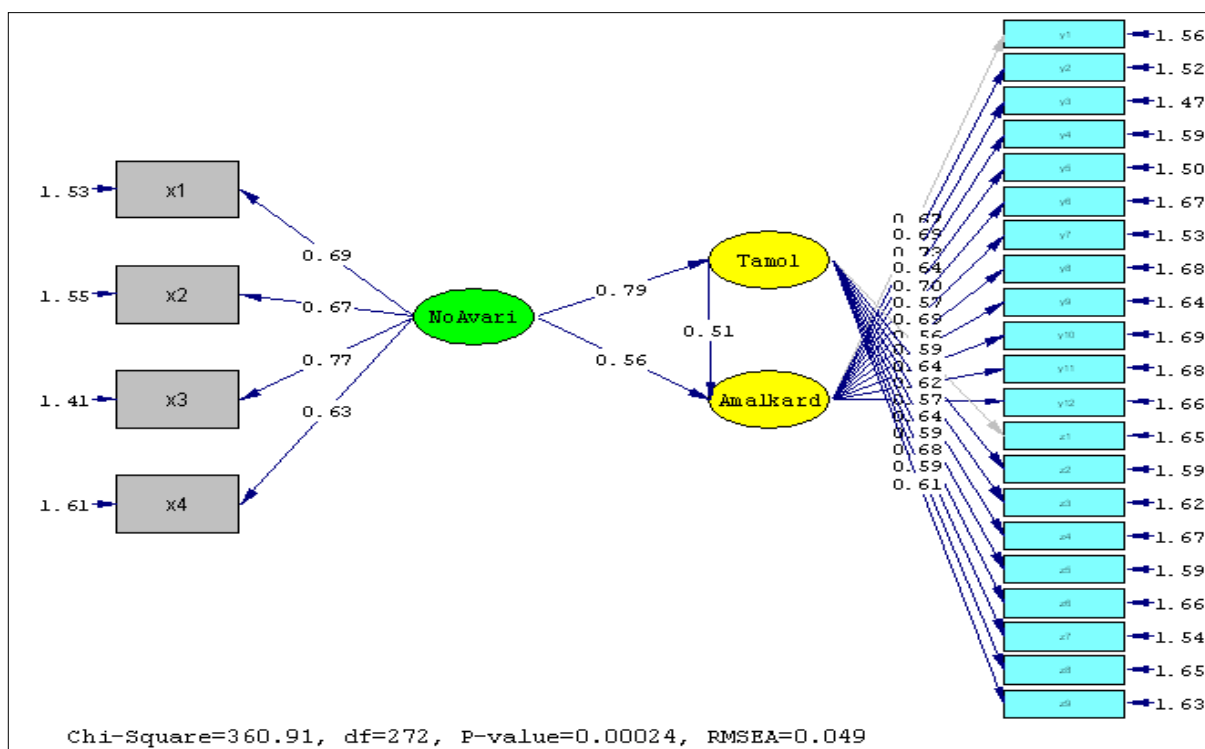
جدول (۲): یافته‌های توصیفی نمونه کمی			
مولفه	ابعاد	فراوانی	درصد
جنسیت	مرد	۱۶۱	۴۱/۹
	زن	۲۲۳	۵۸/۱
سن	کمتر از ۲۵ سال	۲۱	۵/۵
	بین ۲۶ تا ۳۵ سال	۱۴۲	۳۷
	بین ۳۶ تا ۴۵ سال	۱۹۰	۴۹/۵
	بین ۴۶ تا ۵۵ سال	۳۱	۸
	بالای ۵۶ سال	۰	۰
	دیپلم و کمتر	۹	۲/۳
تحصیلات	کاردانی	۵۱	۱۳/۳
	کارشناسی	۲۰۳	۵۲/۹
	کارشناسی ارشد و بالاتر	۱۲۱	۳۱/۵

طور کامل آورده و توضیح داده شده است. همچنین در این شکل میزان خطای مجاز برای هر یک از متغیرهای آشکار مشخص شده است که حاصل محاسبات نرم‌افزاری این متغیرها است و مقادیر آن در خروجی عددی در پیوست آماری این روش (ساختاری) آورده شده است. خروجی کلی نرم افزار برای کل مدل ساختاری تحقیق با در نظر گرفتن تمامی متغیرهای نهفته درون‌زا و برون‌زا در شکل شماره (۱) آورده شده است. در این خروجی، میزان بارهای عاملی برای هر یک از مسیرهای تعریف شده در مدل سازه‌ای تحقیق آورده شده است. نتایج آزمون فرضیه‌ها با استفاده از این خروجی به تفصیل در شکل‌های شماره (۱) و (۲) ارائه می‌شود.

خروجی نرم افزار ضرایب مسیر بررسی رابطه عوامل موثر بر نوگرایی مصرف کننده و عملکرد برندهای آنلاین با توجه به نقش میانجی تعامل مشتریان با برند آنلاین، بین متغیرهای نهفته درون‌زا و برون‌زا را نشان می‌دهد که روابط بین آنها بر اساس فرضیه‌های تحقیق تعریف شده است. در این شکل، متغیرهای مربع، متغیرهای صریح یا سوالات هر متغیر در پرسشنامه تحقیق هستند. از سوی دیگر، متغیرهای شکل دایره، متغیرهای پنهانی هستند که مدل ساختاری تحقیق، قابل تشخیص هستند. اعداد مشخص شده در فلش برای رابطه بین دو متغیر پنهان مدل بر اساس فرضیه‌های تحقیق، همان ضرایب مسیر محاسبه شده برای فرضیه‌های تحقیق است که در جداول زیر به



شکل (۱): خروجی کلی نرم افزار بر اساس شاخص t در مورد مدل ساختاری تحقیق



شکل (۲): خروجی کلی نرم افزار بر اساس شاخص ضریب اثر در مورد مدل ساختاری تحقیق

ارزیابی مدل به وسیله شاخص های برازش؛ شاخص های
برازش مطلق

شاخص های برازش مطلق، شاخص هایی هستند که از یک سو بر اساس تفاوت واریانس ها و کوواریانس های مشاهده شده و از سوی دیگر واریانس ها و کوواریانس های

پیش بینی شده براساس پارامترهای مدل توسعه یافته است. وجود درجه آزادی بالا (نزدیک به مدل استقلال) و کای دو پایین (نزدیک به مدل اشباع) نشان دهنده مقبولیت مدل و مطلوبیت آن است. محاسبات شاخص های برازش مطلق در جدول شماره (۳) آورده شده است.

جدول (۳)، محاسبات شاخص های برازندگی برازش مطلق					
نام فارسی	نام انگلیسی	اختصار	ملاک	مقدار محاسبه شده	تفسیر نتیجه
تقسیم کی دو به درجه آزادی	χ^2/df	CMIN	کمتر از ۳	۱/۳۲	تأیید برازش
سطح معنی داری	p-value	p	کمتر از ۰/۰۵	۰/۰۰۰۲۴	تأیید برازش
جذر میانگین مربعات خطا	Root Mean Squared Residual	RMSEA	۰/۰۸ > RMSEA > ۰/۰۳	۰/۰۴۹	تأیید برازش
شاخص نیکویی برازش	Goodness-of-Fit Index	GFI	بالاتر از ۰/۹	۰/۹۶	تأیید برازش
شاخص نیکویی برازش اصلاحی	Adjusted Goodness-of-Fit Index	AGFI	بالاتر از ۰/۹	۰/۹۲	تأیید برازش

همان گونه که از جدول شماره (۳)، شاخص های برازندگی قابل ملاحظه است، شاخص ها به شرح زیر است، بار عاملی کلیه ابعاد در شکل بالاتر از ۰/۳ است. شاخص برازش χ^2/df معادل ۱/۳۲ شده است، همچنین مقدار آماره RMSEA معادل ۰/۰۴۹ است. ضمناً سطح معنی داری آزمون یا همان p-value که معادل ۰/۰۰۰۲۴ گردیده است. همچنین دیگر شرایط معادله ساختاری نیز برقرار هستند، لذا در سطح ۰/۰۵، فرض صفر را رد و فرض مقابل را مبنی بر اینکه، مدل برازش شده از سطح معنی داری مناسب برخوردار است، پذیرفته می شود.

۵- نتیجه گیری

بر اساس نتایج تحقیق مشخص شد که، تعامل مشتری با برند، رابطه بین عوامل مؤثر بر نوگرایی مصرف کننده با عملکرد برند در فضای آنلاین را میانجیگری می کند. محققان قبلی دی پلسمارک و همکاران (۲۰۱۸)، برایان و همکاران (۲۰۲۰) نتیجه گرفتند، برای ایجاد روابط مؤثر

و بلندمدت، شناخت لایه های زیرین اثرگذار در روابط مصرف کننده ضروری است و درک ناکافی از پیوندهای احساسی تمامی هزینه های بازاریابی صرف شده شرکت را از بین می برد. در نقطه مقابل، تحقیق ها نشان می دهند ایجاد پیوندهای عاطفی، مشتریان را به نگهدارانی سرسخت برای برند مبدل می کند که در صورت جدایی از برند، نگران و مضطرب می شوند و در هنگام در اختیار داشتن برند، مانند مدافعانی سرسخت از برند حمایت می کنند. این پیوندهای احساسی که در طول زمان ایجاد می شوند که همگی با نتایج تحقیق جاری همسو است. مهم ترین علت همسویی آن است که، امروزه بسیاری از محققین و شرکت ها در زمینه تجارت آنلاین به این نتیجه رسیده اند که با ارزش ترین دارایی یک شرکت به منظور بهبود فرآیند بازاریابی، برند و عملکرد برند است. با این وجود، پذیرش برندها در محیط آنلاین یکی از مهم ترین موضوعات در حوزه سیستم های اطلاعاتی و برندها است که کمتر مورد توجه قرار گرفته است. علیرغم مزایای نسبی فراوان تجارت

مشتري به عنوان يك عنصر محتوایی در تبليغات خود اتكا كنند؛ چراكه كه بالذت بخش معرفي كردن تجربه خريد از فروشگاه مي تواند احتمال خريد مشتريان را افزايش داد. ● در كنار بازاریابی سنتی نسبت به توسعه بازاریابی آنلاین اقدام كنند و از آن برای تاثیرگذاري بر مشتريان استفاده شود. به عبارتي، توصيه مي شود كه تمرکز اصلي بازاریابان اين شركت صرفاً بر بازاریابی مستقيم و فروش شخصي توسط كاركنان فروشگاه نباشد و از طريق بازاریابی آنلاین مشتريان را سريع تر و با هزينه كمتر با ويژگي های خدمات اوکالا آشنا کرده و آنها را به خريد ترغيب كنند

● در طراحی وب سايت به ديده گاه های مشتريان توجه خاص شود. همه نقاط قابل بهبود در ويژگي های وب سايت را بررسي كنند، و ويژگي های بهينه ای برای محيط وب سايت خود طراحی كنند. به عبارتي، محيط فروشگاه اينترنتی نیز از انواع بازاریابی استفاده كنند، كه بر ذهنيت مثبت مشتريان از برند اثرگذار باشند و به نوعی تكميل كننده آن باشند. لذا توصيه مي شود، مديران بازاریابی محيط فروشگاه ها را به نوعی طراحی كنند پياده سازی ابعاد بازاریابی آنلاین (يعنی چيدمان، تبليغات شفاهي، تعامل، قيمت و تجربه) در وب سايت به ساده ترين شكل ممكن برای مديران فروش فروشگاه ممكن باشد.

● از ابزار برند به عنوان یکی از عوامل اصلی متقاعد سازی مشتريان استفاده كنند. به زعم محققان، برند مهم ترين عنصری است كه مي تواند تجربه خريد مشتريان را بهبود ببخشد و تجربه بيشتريين قرابت را با بازاریابی آنلاین دارد. به عنوان يك توصيه اجرائی در اين رابطه پيشنهاده مي شود كه مديران بازاریابی شركت در رويه های بازاریابی و تبليغاتی خود به جای فلسفه فروش بر فلسفه بازاریابی اجتماعی تأكيد كنند تا مشتريان احساس كنند كه با خريد از اوکالا سهمی در رفاه کلی جامعه نیز خواهند داشت. به نظر مي رسد كه با توجه به مشكلات بهداشتی ناشی از بحران توليدات نامرغوب و مواد اوليه مصنوعي كه گريبان گیر جامعه است، فلسفه بازاریابی اجتماعی مي تواند به بهترين شكل ممكن نقش تسهيل كننده تصوير برند را نشان دهد و آن را تقويت كند.

آنلاين، در ايران و برخی كشورهای در حال توسعه چندان مورد استقبال قرار نگرفته است. یکی از دلایل عدم خريد آنلاين، عدم پذيرش ريسك است. تحقيقات در سراسر جهان در مورد موانع توسعه تجارت الكترونيك اعتماد را برجسته کرده است. اعتماد به برند كه به عنوان یکی از اجزای شناختی و مؤثر در ارتباط بين مشتري و برند مطرح است، به معنی اعتقاد مشتري به قابل اعتماد بودن، صداقت برند، وفادار بودن برند به وعده های خود است. به اين ترتيب با وجود رقابت بين شركت های ایرانی در جذب مشتري، اين گونه شركت ها همچنان از اينترنت به خصوص شبكه های اجتماعی برای جذب مشتري استفاده نمی كنند.

در سال های اخير، اشتياق به برند یکی از موضوع های مهم در حوزه مطالعات برند بوده است. به عبارتي، زمانی كه يك برند به شكلي عمل مي كند كه توانايی ارضای نيازهای واقعی و ملموس از يك طرف و خواسته های احساسی و غيرملموس از طرف ديگر را داشته باشد، مي تواند احساسی قوی در مصرف كننده ايجاد كند كه مي توان آن را به اشتياق تعبير كرد و اين گونه و رفته رفته مفهوم اشتياق برند در پژوهش های مصرف كننده، به طور گسترده ای در قياس با اشتياق انسانی مفهوم سازی گردیده است. اما تحولات جديد در حوزه فناوری اطلاعات، فضاهای مجازی جذابی مانند شبكه های اجتماعی را ايجاد کرده است كه روز به روز در حال گسترش هستند و عرصه را برای تبليغات كالاهای خدمات توليد كنندگان فراهم مي كنند. با پيشرفت فناوری اينترنت، شركت ها از سايت های شبكه های اجتماعی برای تبليغ و انتشار اطلاعات برند خود استفاده مي كنند. رسانه های اجتماعی ارتباطات سنتی بازاریابی را تغيير داده اند. كاربران اينترنت به تدريج در حال شكل گيري روابط تجاری هستند كه به طور سنتی توسط بازاریابان شكل گرفته است. لذا پيشنهاده مي شود:

● مديران بازاریابی اوکالا در تبليغات رسانه ای خود به نقش تعيين كننده برندینگ و تصوير برند توجه داشته باشند؛ چراكه ويژگي های غيرفيزيكي فروشگاه نیز در كنار ويژگي های فیزیکی آن در متقاعد كردن مشتريان به خريد نقش ايفای مي كند. مديران بازاریابی شركت بر تجربه خريد

منابع

- ۱- رضایی، فائزه؛ حیدرزاده، کامبیز. (۱۳۹۹). بررسی نقش تعدیل کننده تنوع طلبی، نوآوری و مستعد بودن مصرف کننده بر رابطه بین رضایت با وفاداری رفتاری مشتری. مدیریت بازاریابی، ۱۵ (۴۹)، ۱-۲۲.
- ۲- سرگزی، حسین. (۱۳۹۸). ارزیابی مدیریت تبلیغات برند و تاثیرات آن بر تجارت الکترونیک، سومین کنفرانس ملی مطالعات نوین مدیریت و حسابداری در ایران، کرج.
- ۳- سوری، فرید. (۱۳۹۹). تأثیر کیفیت خدمات الکترونیکی و رضایت و اعتماد مشتری در خریدهای آنلاین، کنفرانس بین المللی مدل ها و فن های کمی در مدیریت، قزوین.
- ۴- کاردانی ملکی نژاد، مونا؛ خوراکیان، علیرضا؛ وحیم نیا، فریبرز. (۱۴۰۰). بررسی تاثیر نوجویی ذاتی بر رفتار نوآورانه مصرف کننده با نقش واسطه نوجویی در زمینه خاص (مورد مطالعه: مصرف کنندگان محصولات هوشمند پوشیدنی، مطالعات رفتار مصرف کننده، ۱ (۸)، ۸۵-۱۰۴.
- ۵- منصوری موی، فرشته؛ مرادی، محمد؛ و ملایی، فاطمه. (۱۳۹۷). بررسی تاثیر عملکرد بازاریابی خدمات بر تبلیغات شفاهی: نقش ارزش ادراک شده مشتری، تجربه مشتری، واکنش احساسی و وفاداری برند. مطالعات مدیریت گردشگری (مطالعات جهانگردی)، ۳۹. ۷۲-۴۹.
- ۶- نوایی، حمیدرضا. (۱۳۹۶). مروری بر عوامل موثر بازاریابی رسانه اجتماعی بر بازاریابی اجتماعی، پنجمین همایش دانشی تحقیقی یافته های نوین علوم مدیریت، کارآفرینی و آموزش ایران، تهران، انجمن توسعه و ترویج علوم و تکنیک های بنیادین.
- ۷- وظیفه دوست، حسین؛ معماریان، شیما. (۱۳۹۵). رابطه رفتار اخلاقی فروشنده با رضایت، اعتماد و وفاداری بیمه گذاران در بیمه های عمر، نشریه پژوهشکده بیمه، ۱ (۱)، ۱۵۸-۱۲۷.
- ۸- یزدانی، سید محمد علی؛ و راکی زاده، محمد. (۱۳۹۶). تأثیر برندها در بین مشتریان فروشگاه های مجازی پوشاک در گرایش آنان به مد، کنفرانس مدیریت و علوم رفتاری، تهران.

- 9-Ahmadi, A. (2019). Thai Airways: key influencing factors on customers' word of mouth. *International Journal of Quality & Reliability Management*, 36(1), 40-57.
- 10-Allen, J., Jochims, B., & Wu, S. (Eds.). (2022). *Celebrating the Past and Future of Marketing and Discovery with Social Impact: 2021 AMS Virtual Annual Conference and World Marketing Congress*. Springer Nature.
- 11-Baldus, B. J., Voorhees, C., & Calantone, R. (2015). Online brand community engagement: Scale development and validation. *Journal of business research*, 68(5), 978-985.
- 12-Bhuvaneswari, V., & Krishnan, J. (2015). A review of literature on impulse buying behaviour of consumers in brick & mortar and click only stores. *International journal of management research and social science*, 2(3), 84-90.
- 13-De Pelsmacker, P., Van Tilburg, S., & Holthof, C. (2018). Digital marketing strategies, online reviews and hotel performance. *International Journal of Hospitality Management*, 72, 47-55.
- 14-Hamilton, R. G. (2019). Assessment of human allergic diseases. In *Clinical immunology* (pp. 1283-1295). Elsevier.
- 15-Jiang, M., Yang, J., Joo, E., & Kim, T. (2022). The Effect of Ad Authenticity on Advertising Value and Consumer Engagement: A Case Study of COVID-19 Video Ads. *Journal of Interactive Advertising*, 22(2), 178-186.
- 16-Jin, B., & Kang, J. H. (2011). Purchase intention of Chinese consumers toward a US apparel brand: A test of a composite behavior intention model. *Journal of consumer marketing*, 28(3), 187-199.
- 17-Lebedeva, M. G., Lupo, A. R., Henson, C. B., Solovyov, A. B., Chendev, Y. G., & Market, P. S. (2018). A comparison of bioclimatic potential in two global regions during the late twentieth century and early twenty-first century. *International journal of biometeorology*, 62, 609-620.
- 18-Mitchell, V. W., & Balabanis, G. (2021). The role of brand strength, type, image and product-category fit in retail brand collaborations. *Journal of Retailing and Consumer Services*, 60, 102445.
- 19-Pantano, E. (2021). When a luxury brand bursts: Modelling the social media viral effects of negative stereotypes adoption leading to brand hate. *Journal of Business Research*, 123, 117-125.
- 20-Pantano, E., Dennis, C., & De Pietro, M. (2021). Shopping centers revisited: the interplay between consumers' spontaneous online communications and retail planning. *Journal of Retailing and Consumer Services*, 61, 102576.
- 21-Vander Schee, B. A., Peltier, J., & Dahl, A. J. (2020). Antecedent consumer factors, consequential branding outcomes and measures of online consumer engagement: current research and future directions. *Journal of Research in Interactive Marketing*, 14(2), 239-268.
- 22-Vickers, N. J. (2017). Animal communication: when i'm calling you, will you answer too?. *Current biology*, 27(14), R713-R715.
- 23-Yuan, D., Lin, Z., Filieri, R., Liu, R., & Zheng, M. (2020). Managing the product-harm crisis in the digital era: The role of consumer online brand community engagement. *Journal of Business Research*, 115, 38-47.

©Authors, Published by Journal of Intelligent Knowledge Exploration and Processing. This is an open-access paper distributed under the CC BY (license <http://creativecommons.org/licenses/by/4.0/>).



مقاله پژوهشی

تشخیص لبه تصاویر دیجیتالی با استفاده از خوشه‌بندی فازی بهینه شده

Doi: 10.30508/KDIP.2023.368960.1054

محب علی صوفی (نویسنده مسئول)^۱ | علیرضا نصیری^۲

۱- دانش آموخته کارشناسی ارشد هوش مصنوعی و ریاتیکز، دانشگاه هرمزگان، هرمزگان، ایران

۲- استادیار گروه برق، دانشگاه هرمزگان، هرمزگان، ایران

تاریخ دریافت: ۱۴۰۱/۰۸/۲۲

تاریخ پذیرش: ۱۴۰۱/۱۰/۱۹

صفحه: ۵۵ - ۵۰

چکیده:

یکی از مهم‌ترین بخش‌های تصویر لبه‌ها می‌باشند. به مرز بین دو ناحیه از تصویر که دارای تفاوت قابل توجه در شدت روشنایی، رنگ و یا بافت باشد، لبه گفته می‌شود. تشخیص لبه یک مرحله اساسی در تقسیم‌بندی تصویر و یکی از مهم‌ترین مراحل برای تشخیص ویژگی اشیاء در یک تصویر می‌باشد. اخیراً پژوهش‌های زیادی در زمینه تشخیص لبه انجام شده است که هر کدام سعی در تشخیص بهتر لبه‌های تصویر داشته‌اند. در این پژوهش برای تشخیص لبه تصاویر دیجیتالی، از تئوری فازی و مفهوم خوشه‌بندی فازی استفاده شده است، در ابتدا پیکسل‌های تصویر بر اساس میزان سطح خاکستری آنها به خوشه‌هایی تقسیم شده و سپس با تعیین فیلترهای مناسب اقدام به تشخیص لبه تصویر شده است. روش پیشنهاد شده لبه‌های تصویر را به خوبی تشخیص داده، همچنین در تشخیص لبه‌هایی از تصویر که به علت سطح نور نامناسب و یا نویزی بودن امکان تشخیص صحیح آنها وجود ندارد عملکرد چشمگیری از خود نشان داده است. نتایج این پژوهش نشان داد؛ روش پیشنهادی در اکثر موارد نتایج بهتری نسبت به روش‌های معمول ارائه داده است و این به علت عملکرد خوب آن در تشخیص جزئیات تصویر مانند؛ ضخامت و قوی و ضعیف بودن لبه‌ها می‌باشد به خصوص در تصاویری که دارای وضوح کمتر و یا کانتیر پایین‌تری است، عملکرد خیلی خوبی داشت.

کلمات کلیدی: پردازش تصویر، تشخیص لبه، تئوری فازی، خوشه‌بندی فازی.

۱- مقدمه

تصویر یکی از بزرگ‌ترین منابع اطلاعاتی فعالیت‌های هوشمندانه بشر است. با گذشت زمان نیاز به پردازش تصویر نیز روز به روز افزایش پیدا کرده است (کومار و راهجا، ۲۰۲۰). تشخیص لبه یکی از عملیات‌های مورد استفاده در تجزیه و تحلیل تصویر و همچنین پیش پردازش لازم در تقسیم‌بندی تصویر است. یک لبه مرز میان دو ناحیه تصویر است. عملیات تشخیص لبه در حوزه‌های مختلف پردازش تصویر مانند؛ استخراج ویژگی‌ها، تشخیص ویژگی‌ها و تشخیص الگو و غیره ضروری است (سینگه، شکار، سینگه و چوهان، ۲۰۱۴). تمام روش‌های تشخیص لبه تحت دو گروه گرادیان و لاپلاس طبقه‌بندی می‌شوند (بستان، بخاری، وبروئل، ۲۰۱۷؛ کروویلا، سوکوماران، سنکار، و جوی، ۲۰۱۶). در بین روش‌های تشخیص لبه، عملگر لاپلاسن، روبرتز، سوبل، پرویت، کنی و روش‌های منطق فازی دارای بیشترین استفاده هستند، وجود نویز و مات‌شدگی تصویر از جمله مشکلاتی هستند که در این تکنیک‌ها با آن مواجه می‌شوید (راهجا و کومار، ۲۰۲۱). انعطاف‌پذیری و پتانسیل ذاتی سیستم‌های فازی از جمله مقاومت در برابر نویز و سهولت ادغام با دستگاه‌های نانومتریک برای پیاده‌سازی‌های عملی، آن‌ها را به بستری مناسب برای کاربردهای مختلف تخصصی دانش بنیان تبدیل کرده است (بزرگمهر، معیری، ناوی، و باقرزاده، ۲۰۱۷). در پژوهش‌رن و همکاران که با هدف ارائه یک رویکرد فازی جدید انجام شده، برای تشخیص لبه تصاویر دیجیتالی از روش خوشه‌بندی فازی استفاده شده است. خوشه‌بندی یکی از شاخه‌های یادگیری بدون نظارت می‌باشد و فرآیند خودکاری است که در طی آن نمونه‌ها به دسته‌هایی که اعضای آن مشابه یکدیگر می‌باشند تقسیم می‌شوند (رن، وانگ، فنگ، اکسو، لیو، و دینگ، ۲۰۱۹).

۲- مبانی نظری

یکی از روش‌های آشکارسازی لبه، استفاده از عملگرهای کلاسیک مانند فیلترهای خطی می‌باشد. این روش‌ها اگر چه دارای قدرت کمتری در آشکارسازی لبه هستند ولی به دلیل سهولت پیاده‌سازی بیشتر مورد استفاده قرار می‌گیرند. در پژوهشی یک روش فازی برای حذف نویز و ضربه در تصاویر دیجیتال بررسی شده است، روش پیشنهادی از تابع عضویت مثلثی از نوع ممدانی و پنجره 2×2 استفاده کرده است، در نهایت نتیجه لبه‌یابی به دست آمده با عملگرکانی^۷ و سوبل^۸ مقایسه شده است (بزرگمهر، جوک، موآجری، ناوی، باقرزاده، ۲۰۲۰). باهاگمباتی و همکاران یک روش جدید بر پایه استراتژی استدلال منطق فازی برای تشخیص لبه در تصاویر دیجیتال بدون تعیین مقدار آستانه پیشنهاد کرده‌اند، روش پیشنهاد شده با بخش‌بندی تصویر به نواحی مختلف، با استفاده از ماتریس باینری 3×3 شروع می‌شود، به طوری که پیکسل‌های لبه به محدوده یا از مقادیر متمایز از یکدیگر نگاشته می‌شود. قابلیت اطمینان نتایج روش پیشنهاد در این مقاله شده برای تصاویر گرفته شده متفاوت، با عملگر سوبل خطی مورد مقایسه قرار می‌گیرد، این روش باعث تأثیر پایدار در همواری و صاف بودن خطوط برای خط‌های مستقیم و گرد شدن برای خط‌های منحنی می‌شود. در این حالت گوشه‌های تصویر تیزتر و به راحتی قابل تعریف می‌باشند. از جمله معایب روش مذکور، این است، هنگامی که عملگر سوبل بر روی این تصویر اعمال می‌شود، لبه‌های گسسته یا در سمت چپ ظاهر خواهد شد، که در نتیجه باید در انتخاب قوانین فازی خاص که باعث اجتناب از دوبر شدن لبه‌ها و به دست آوردن یک تصویر با لبه‌های منفرد شود دقت بسیاری نمود (باهاگمباتی و داس، ۲۰۱۲). در پژوهشی نتایج مطالعه یک روش جدید فازی برای تشخیص لبه تصویر بر اساس ترانزیستور اثر میدان

1- Kumar, & Raheja

2- Singh, Shekhar, Singh, & Chauhan

3- Batan, Bukhari, & Breuel,

4- Kuruvilla, Sukumaran, Sankar, & Joy

5- Bozorgmehr, Moaiyeri, Navi, & Bagherzadeh

6- Ren, Wang, Feng, Xu, Liu, & Ding

7- Canny

8- Sobel

9- Bozorgmehr, Jooq, Moaiyeri, Navi, & Bagherzadeh

10- Bhagabati, & Das

که حاوی تعداد زیادی باند با شکاف بسیار باریک در حوزه طیفی است، با مشکلاتی مواجه هستند. در این پژوهش، یک روش جدید بر مبنای منطق فازی برای تشخیص لبه در تصاویر دیجیتال که توانایی یافتن آستانه را ندارند، پیشنهاد شده است. این روش از طریق تقسیم بندی تصاویر به بخش هایی با استفاده از ماتریس های 3×3 دوتایی آغاز می شود. این روش پیشنهادی با نتایج حاصل از روش های PSO و شبکه عصبی مقایسه شده است. این تکنیک پیشنهادی یک اثر دایمی در هموارسازی خطوط و صافی برای خطوط مستقیم و گردی خوب برای خطوط منحنی و وضوح در منحنی ها ارائه می کند (دیوا و پراکاش^۱، ۲۰۱۹).

۳- روش تحقیق

الگوریتم با قسمت کردن پیکسل های تصویر HI به ناحیه های (ماسک های) 3×3 شروع می شود و سپس تصویر به وسیله سیستم استنتاج فازی به محدوده ای از مقادیر نگاشت داده می شود و تشخیص لبه انجام می شود (خیار و تکار^۲، ۲۰۱۲). مراحل اصلی در تشخیص لبه به روش فازی عبارتند از: فازی سازی تصاویر، اصلاح مقادیر عضویت و غیر فازی سازی در صورت لزوم. با توجه به هشت بیتی بودن تصویر ورودی و تصویر خروجی ای که بعد از غیر فازی سازی بدست می آید، سطوح خاکستری همیشه بین مقادیر ۰ تا ۲۵۵ می باشند. مجموعه های فازی برای نشان دادن شدت هر یک از متغیرها تولید می شوند. این مجموعه ها به متغیرهای زبانی رنگ مشکی، رنگ سفید و لبه وابسته اند، همچنین تابع عضویت در نظر گرفته شده برای مجموعه های فازی ورودی و خروجی به صورت مثلثی می باشد. در مرحله استنتاج نیز از توابع Min و Max برای عملیات And و Or استفاده می شود، روال غیر فازی سازی و استنتاج نیز روش ممدانی می باشد (اوروجو، مسکلانس، دامسویوس، و ووی^۳، ۲۰۲۰). مراحل فازی و غیر فازی کردن به دلیل انجام نشدن پردازش سخت افزار فازی است. بنابراین کد کردن داده تصویر و دیکد کردن نتایج مرحله

نانولوله کربنی (CNTFET) ارائه می کند. طراحی پیشنهاد شده در این مقاله از یک قاعده ساده برای تشخیص لبه تصویر دیجیتال استفاده می کند. برای بررسی عملکرد و رویکرد فازی پیشنهادی در این مقاله نتایج با اپراتورهای تشخیص لبه MATLAB مانند Canny, Sobel و Prewitt مقایسه شده اند و نتایج شبیه سازی، نشان داد که رویکرد فازی پیشنهادی می تواند از طریق گیت های منطقی در مقیاس نانو پیاده سازی شود و دقت معقولی را ارائه نماید و عنصر اصلی سیستم پیشنهاد شده که با CNTFETs اجرا می شود، مساحت چیدمان $1/0$ را اشغال کرده و 535 نانو وات توان را مصرف می کند. نتایج این مقاله همچنین بر فواید گیت- همه- طرفه (GAA) به عنوان یک کاندیدای بالقوه برای کاربردهای پردازش تصویر فازی در مقیاس نانو تاکید می کند (بزرگمهر و همکاران، ۲۰۲۰). در پژوهش دیگری، فلورس ویدال و همکاران که عملیات تشخیص لبه را با متد خوشه بندی فازی انجام داده، مشاهده می شود؛ به طور سنتی، فرآیند تشخیص لبه به یک مرحله نهایی نیاز دارد که به عنوان مقیاس بندی شناخته می شود. این کار برای تصمیم گیری، پیکسل به پیکسل برای مشخص شدن اینکه این کار به عنوان لبه نهایی انتخاب شود یا نه، انجام می شود. این مساله را می توان به عنوان یک روش ارزیابی محلی در نظر گرفت که مشکلات زیادی در عمل ایجاد می کند، زیرا پیکسل های حاشیه نباید به عنوان مستقل در نظر گرفته شوند (فلوراس ویدال، اولاسو، گومز، و گیودا^۴، ۲۰۱۹). نتایج پژوهش دیوا و همکارش، نشان داد که تشخیص لبه نقش مهمی در زمینه پردازش تصویر دارد. تکنیک های مختلف تشخیص لبه مانند سوبل، PSO، Prewitt، لاپلاسن و لاپلاسن گاوسی به دست آمده است. این تکنیک ها محدودیت هایی مانند ضخامت لبه ثابت را مصرف می کنند و برخی پارامترها مانند: آستانه برای پیاده سازی، مشکل ساز است. اگر چه بسیاری از روش های تشخیص لبه برای تصاویر در مقیاس خاکستری، رنگی و چند طیفی پیشنهاد شده اند، اما هنوز هنگام استخراج ویژگی های لبه از تصاویر فراطیفی (HSIs)

1- Flores-Vidal, Olaso, Gómez, & Guada

2- Dhivya, & Prakash

3- Khaire, & Thakur

4- Orujov, Maskelinas, Damaševičius, & Wei

$$\theta = \cos^{-1} \left\{ \frac{0.5[(R-G) + (R-B)]}{[(R-G)^2 + (R-B)(G-B)]^{0.5}} \right\} \quad (2-3)$$

مولفه‌ی اشباع براساس معادله‌ی (۳-۳) بدست می‌آید:

$$S = 1 - \frac{3}{(R+G+B)} \min(R, G, B) \quad (3-3)$$

سرانجام مولفه‌ی شدت نیز براساس معادله‌ی (۳-۴) بدست می‌آید:

$$I = \frac{1}{3} (R + G + B) \quad (4-3)$$

فرض می‌شود که مقادیر RGB به بازه‌ی [۰, ۱] نرمال‌سازی شدند و زاویه θ نسبت به محور قرمز در فضای HSI محاسبه می‌شود. پرده رنگ H می‌تواند به بازه‌ی [۰, ۱] نرمال شود، که برای این کار باید تمام مقادیر حاصل از معادله ۳-۳ بر ۳۶۰ درجه تقسیم گردد. اگر مقادیر RGB در فاصله [۰, ۱] باشند، دو مولفه دیگر HSI نیز در این بازه بوده‌اند.

انسان ابتدا عکس‌ها را بر اساس دامنه‌ی طول موج نوری که از یک شیء منعکس می‌شود به عبارت دیگر رنگ و حجم نوری که از آن جسم منعکس می‌شود، یعنی شدت می‌بیند. با استفاده از این واقعیت مولفه اشباع از تصویر HSI حذف شده و شدت و رنگ برای تشکیل فرم جدیدی از عکس‌های HI اضافه می‌شوند. مقادیر پیکسل‌ها در محدوده‌ی [۰, ۱] به [۰, ۲۵۵] برای تسهیل در فهم محاسبات تغییر می‌یابند. عکس‌های فراهم شده‌ی HI شبیه رنگ خاکستری با مقادیر پیکسل صفر تا ۲۵۵ به نظر می‌رسد (بیسواس، و هزرا^۳، ۲۰۱۸؛ کالرا، و چوهکار^۴، ۲۰۱۶).

فازی کردن، محاسبه‌ی درجه‌ی عضویت و خوشه‌بندی پیکسل‌ها

این عکس‌های HI به‌عنوان مجموعه‌ی فازی در نظر گرفته شده است. درجه عضویت فازی عناصر پیکسل بر اساس مقادیر خاکستری HI تعریف می‌شود به طوری که بیشترین تعداد پیکسل‌ها که مقادیر خاکستری ثابتی دارند، دارای بالاترین درجه‌ی عضویت می‌باشند، به طور مشابه مجموعه‌ی دومین بیشترین پیکسل‌ها که دارای

است که امکان پردازش تصویر را با تکنیک‌های فازی فراهم می‌سازد، قدرت اصلی در مرحله میانی یعنی اصلاح مقادیر عضویت فازی می‌باشد. سپس داده‌های تصویر از صفحه سطوح خاکستری به صفحه عضویت فازی تبدیل می‌شود. تکنیک‌های فازی مناسب مقادیر عضویت فازی را اصلاح کرده و تصویر را خوشه‌بندی می‌کند (وانگ، ژانگ، و ویی^۱، ۲۰۱۹).

بنابراین لبه‌یابی تصویر پیشنهادی براساس سیستم فازی شامل ۵ مرحله اصلی به شرح زیر است:

- تبدیل تصویر رنگی RGB به تصویر HSI و HI
- فازی کردن پیکسل‌ها و محاسبه مقادیر درجه عضویت فازی
- خوشه‌بندی تصویر و تعریف قوانین فازی
- اعمال سیستم استنتاج ممدانی
- غیرفازی‌سازی کردن خروجی

تبدیل تصویر رنگی RGB به تصویر HSI و HI

مدل HSI ابزار ایده‌آلی برای تولید الگوریتم‌های پردازش تصویر براساس توصیف‌های رنگی است که برای انسان، طبیعی و شهودی است. به طور خلاصه می‌توان گفت که RGB برای تولید رنگ تصویر مناسب است (مثل تصویربرداری توسط دوربین خانگی یا نمایش تصویر در صفحه مانیتور رنگی)، اما استفاده از آن برای توصیف رنگ، خیلی محدود است. مطالبی که در ادامه می‌آید، روش آسانی را برای انجام این کار فراهم می‌سازد (کیو و همکاران^۲، ۲۰۲۱). محاسبات از RGB به HSI و برعکس، برحسب بیت‌ها انجام می‌گیرد. برای سهولت، وابستگی (x, y) در معادلات تبدیلی در نظر گرفته نشده است. اگر تصویری با فرمت RGB داشته باشیم، مولفه H هر پیکسل RGB با استفاده از معادله (۳-۱) به دست می‌آید.

$$H = \begin{cases} \theta & \text{if } B \leq G \\ 360 - \theta & \text{if } B > G \end{cases} \quad (1-3)$$

بطوری که مقدار θ از رابطه‌ی (۲-۳) بدست می‌آید

1- Wang, Zhang, & Wei

2- Qi et al

3- Biswas, & Hazra

4- Kalra, & Chhokar

ماسک‌های G_x و G_y به منظور تشخیص لبه در جهت افقی و عمودی به کار می‌رود.

G_x	G_y		
0	1	0	-1
-1			

شکل (۲): ماسک‌های لبه

دامنه‌ی پیکسل‌های لبه با استفاده از معادله‌ی (۳-۶) محاسبه می‌شود:

$$G = \sqrt{G_x^2 + G_y^2} \quad (3-6)$$

این ماسک‌های 3×3 به محاسبات کمتری برای تشخیص لبه در مقایسه با ماسک‌های 3×3 دیگر بکاررفته در Sobel و Prewitt نیاز دارد. همچنین اثر Blurring به هنگام تشخیص لبه کاهش می‌یابد. عموماً عکس‌های واقعی شامل هر دو لبه قوی و ضعیف می‌باشد. اینجا دو آستانه برای لبه‌ها تنظیم شده که شامل آستانه بالا و آستانه پایین می‌شود. لبه‌های بالاتر از آستانه بالا، لبه‌های قوی و لبه‌های پایین‌تر از آستانه پایین، لبه‌های ضعیف می‌باشند. مقادیر آستانه بالا برای لبه‌های قوی 0.3 و برای لبه‌های ضعیف مقدار آستانه پایین برابر 0.4 ضرب در آستانه بالا می‌باشد.

تعریف قوانین فازی

ویرایش‌گر قاعده فازی شامل عملکرد تعیین قوانین به تابع عضویت است. قوانین می‌توانند بر اساس انواع ورودی تابع عضویت تغییر کنند. به طور کلی ماسک‌های 3×3 بر اساس ۲۸ قانون عمل می‌کنند (بیسواس، و هزرا، ۲۰۱۸). همان‌طور که در شکل شماره (۳) مشخص است این ماسک‌ها با توجه به حالت سیاه (b) یا سفید (w) بودن پیکسل‌های اطراف یک پیکسل، لبه بودن آن پیکسل را مشخص می‌کنند.

$$\mu_{A \rightarrow B}(x, y) = \mu_R(x, y) = \mu_A(x) \wedge \mu_B(y) \\ = \min[\mu_A(x), \mu_B(y)].$$

مقادیر ثابت خاکستری می‌باشند، درجه عضویت بعدی را داراست. به عنوان مثال کمتر از یک. هر پیکسل در عکس مقادیر عضویت خودش را بر اساس تعداد پیکسل‌های با همان درجه خاکستری نگه می‌دارد. حال یک پیکسل در این عکس فازی سه مشخصه دارد:

۱- مختصات فضایی

۲- مقدار پیکسل (مقدار خاکستری)

۳- درجه عضویت فازی (مقدار عضویت)

مجموعه فازی F از عکس‌ها بر اساس معادله‌ی ۵-۳ تعریف می‌شود: -

$$F = \{(x, \mu_F(x), x \in X)\} \quad (3-5)$$

که $\mu_F(x)$ مقدار درجه عضویت (پیکسل) عنصر x در (عکس) مجموعه فازی F را مشخص می‌کند. که میزان مقادیر خاکستری ثابت برای هر پیکسل نشان‌دهنده درجه عضویت آن پیکسل و میزان تعلق آن به هر خوشه می‌باشد.

اشکال متفاوتی از توابع عضویت از قبیل: مثلثی، دوزنقه‌ای، تک‌های خطی، گوسین، زنگوله و غیره وجود دارد. در بسیاری از مقالات تابع عضویت مثلثی به علت دارا بودن فرمول‌های انعطاف‌پذیر برای محاسبه استفاده می‌شود. پارامترهای تابع عضویت مثلثی به عنوان a, b, c است که b نقطه اوج آن است. در روش پیشنهادی در این مقاله پیکسل‌های ورودی به مجموعه سیاه، سفید و پیکسل خروجی به سه مجموعه فازی سیاه (B)، لبه (E) و سفید (W) تقسیم شده و از ماسک‌های 3×3 استفاده شده است.

شکل (۱): حالت کلی ماسک‌های 3×3		
$(1-j, 1-i)$	$(j, 1-i)$	$(1+j, 1-i)$
$(1-j, j)$	(j, j)	$(1+j, j)$
$(1-j, 1+i)$	$(j, 1+i)$	$(1+j, 1+i)$

یک اپراتور گردایان 3×3 در جهت افقی و عمودی در شکل شماره (۱) نشان داده شده است. این ماسک‌ها بر روی عکس‌های تقسیم شده جزئی (HI) مرحله‌ی قبل و

Rule 1	Rule 2	Rule 3	Rule 4	Rule 13	Rule 14	Rule 15	Rule 16
Rule 5	Rule 6	Rule 7	Rule 8	Rule 21	Rule 22	Rule 23	Rule 24
Rule 9	Rule 10	Rule 11	Rule 12				

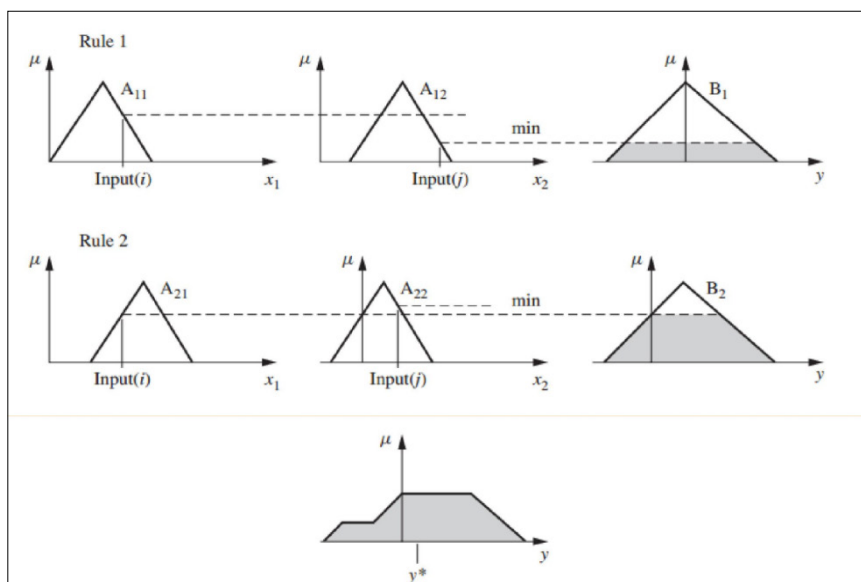
شکل (۳): ۲۸ قانون فازی سیستم پیشنهادی

سیستم استنتاج ممدانی

فرم کلی قوانین در سیستم استنتاج ممدانی در شکل شماره (۴) نشان داده شده است، نتایج هر قانون که با مقدارهای B_1 و B_2 مشخص می‌باشند. همان طور که مشخص است پس از تجمیع نتایج حاصل قوانین

براساس معادله‌ی (۳-۷)، عمل غیرفازی‌سازی بر روی تابع عضویت خروجی انجام می‌گیرد و نتایج به صورت عددی بدست می‌آید (دریانکو، هلندوم، و رینفرانک، ۲۰۱۳).

$$\mu_{A \rightarrow B}(x, y) = \mu_R(x, y) = \mu_A(x) \wedge \mu_B(y) \quad (3-7) \\ = \min[\mu_A(x), \mu_B(y)].$$



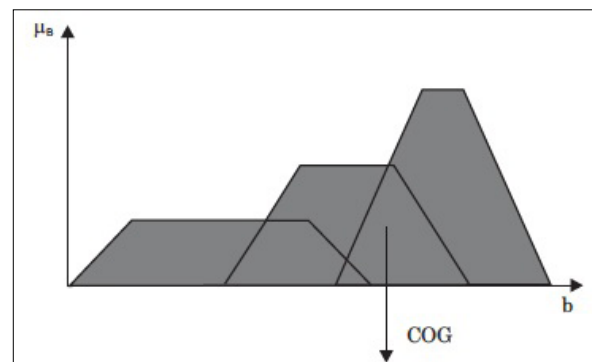
شکل (۴): فرم کلی اجرای قوانین در سیستم استنتاج ممدانی

غيرفازی کردن خروجی

غيرفازی کردن با روش‌های مختلفی انجام می‌گیرد که متداول‌ترین آن روش مرکز ثقل^۱ (COG) می‌باشد. غيرفازی ساز مرکز ثقل نقطه را بعنوان مرکز ناحیه‌ای که بوسیله تابع عضویت پوشش داده شده تعریف می‌کند که در رابطه (۳-۸) مشخص شده است.

$$COG(A) = \frac{\int \mu_A(z).z dz}{\int \mu_A(z).dz} \quad (۸-۳)$$

نمایش گرافیکی این غيرفازی ساز در شکل شماره (۵) نشان داده شده است.



شکل (۵): مرکز ثقل در سیستم فازی

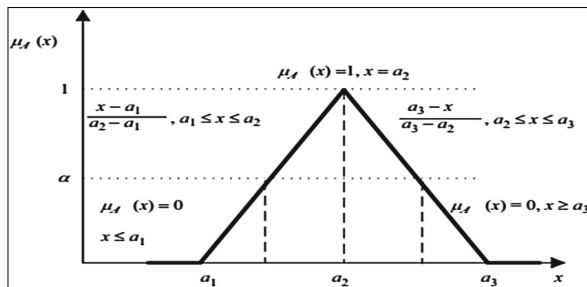
۴- یافته‌های تحقیق

معرفی مجموعه داده

در این پژوهش از مجموعه داده^۲ BSD استفاده شده است. این پایگاه داده حاوی قسمت‌هایی از Ground Truth می‌باشد که این تصاویر توسط انسان حاوی طیف وسیعی از صحنه‌های طبیعی تولید شده است. معیارهایی که برای انتخاب گرفته شده بر این اساس بوده که حداقل در هر تصویر، منظره طبیعی یک شی باید قابل شناسایی باشد.

تابع عضویت فازی

در منطق فازی این تابع نشان دهنده درجه حقیقت بعنوان بسطی از ارزیابی است. در این پژوهش از تابع عضویت مثلثی به دلیل انعطاف پذیری بیشتر استفاده شده است. پارامترهای تابع عضویت مثلثی به عنوان a, b, c است که b نقطه اوج آن می‌باشد.



شکل (۶): فرم کلی تابع عضویت مثلثی

تنظیمات و نتایج آزمایشات

این روش پیشنهادی با استفاده از متلب ۲۰۱۸ شبیه‌سازی شده است. تصاویر BSD و Ground Truth مربوط در آزمایشات استفاده شد. پارامترهای بهبود استفاده شده PSNR و PR (نسبت لبه‌های درست به غلط) هستند.

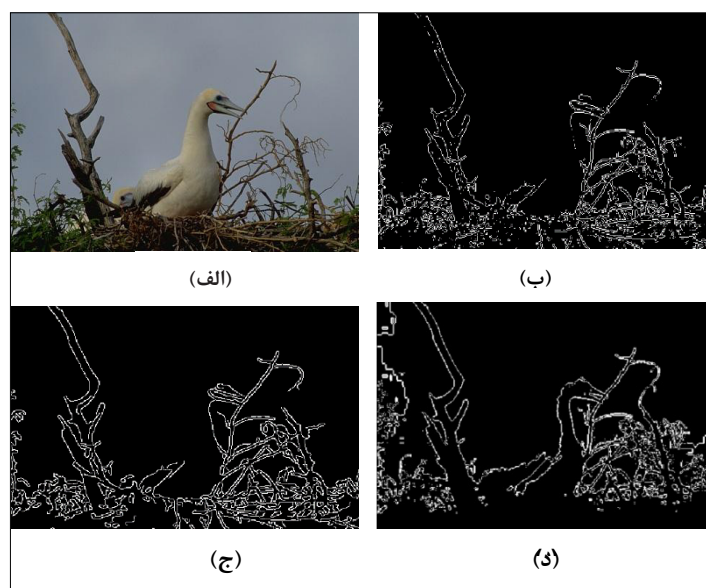
$$PSNR(f, g) = 10 \log_{10} (255^2 / MSE(f, g))$$

$$MSE(f, g) = \frac{1}{MN} \sum_{i=1}^M \sum_{j=1}^N (f_{ij} - g_{ij})^2 \quad (۲-۴)$$

لبه‌یابی تصاویر در روش پیشنهادی

پس از نوشتن کد برنامه در نرم افزار مطلب عکس‌های انتخاب شده به برنامه داده شدند، در شکل شماره (۷)، یکی از تصاویر لبه‌یابی شده مشاهده می‌شود.

1- Center of gravity
2- Berkeley Segmentation Dataset



شکل (۷): تصویر پرند لبه یابی شده به روش پیشنهادی. (الف) تصویر پیش فرض، (ب) تصویر لبه یابی شده با روش کنی، (ج) تصویر لبه یابی شده با روش سوئل، (د) تصویر لبه یابی شده به روش پیشنهادی

در روش پیشنهادی تشخیص لبه ها به وضوح قابل مشاهده هستند و از کیفیت و وضوح بیشتری نسبت به روش کنی و سوئل برخوردار است. این روش تصاویر واضح تری نسبت به روش های مرسوم سوئل و کنی ارائه داده است که این به علت قدرت این روش در تشخیص لبه های ضعیف می باشد. نتایج مربوط به عملکرد تعداد ۳۰ عکس که با روش پیشنهاد شده و روش کنی لبه یابی و مقایسه شده اند در جدول شماره (۲) مشاهده می شود، بیشترین مقادیر PR و PSNR در روش پیشنهاد شده مربوط به عکس ۱ (۱۳۵۰۶۹) به ترتیب مقادیر ۰/۴۱۹ و ۹۹۳۱/۲۳ می باشد. میانگین PR مجموعه تصاویر ارزیابی شده به روش پیشنهاد شده ۷۰۰/۱۴ می باشد که نسبت به میانگین PR در تصاویر بدست آمده با روش کنی به مراتب بهتر و دقیق تر است. میانگین PSNR تصاویر ارزیابی شده به روش پیشنهاد شده ۳۵۸/۲۰ می باشد که نسبت به میانگین PSNR مربوط به تصاویر مربوط به روش کنی نسبتاً بهتر است.

همان طور که در شکل شماره (۷) مشاهده می شود، در روش پیشنهادی لبه های تشخیص داده شده به وضوح قابل مشاهده بوده و از کیفیت و وضوح بیشتری نسبت به روش کنی و سوئل برخوردار است. در جدول شماره (۱) مقادیر PR و PSNR، شکل شماره (۷)، روش پیشنهادی با روش های کنی و سوئل مقایسه شده است، چنانچه مشاهده می شود مقدار PR در روش پیشنهادی بسیار بهتر از مقادیر PR در روش سوئل و کنی است ولی نسبت PSNR روش پیشنهادی با روش کنی و سوئل دارای اختلاف نسبتاً کمی است.

جدول (۱): مقادیر PR و PSNR شکل (۷) با روش پیشنهادی و روش های کنی و سوئل		
روش / خطا	PR	PSNR
سوئل	۱/۳۸	۴۳/۳۵
کنی	۲/۶۳	۴۳/۳۸
روش پیشنهادی	۳/۷۹	۴۳/۳۷

جدول (۲): مقایسه نتایج ارزیابی تصاویر روش پیشنهاد شده با روش کنی

	BSD Image	Canny (T=0.3, σ=1)		Proposed (T=0.3)	
	.No	(PSNR(dB	PR	(PSNR(dB	PR
1	135069	9735/23	1513/14	9931/23	0419/28
2	176039	4721/23	9261/6	487/23	113/10
3	15088	3726/23	8939/8	3953/23	7663/17
4	12074	2335/23	2393/7	246/23	8485/10
5	210088	0153/23	7167/8	0247/23	2036/13
6	28075	9767/22	2033/8	981/22	1661/8
7	108073	5102/21	8069/5	5224/21	3232/10
8	3096	4189/21	5835/5	4337/21	5958/9
9	134052	2264/21	4736/6	2345/21	0943/11
10	189080	9795/20	993/8	9949/20	4702/12
11	189011	6882/20	3055/6	6976/20	5861/10
12	253036	4877/20	6724/12	4997/20	8664/21
13	8068	1962/20	9881/4	1987/20	951/4
14	310007	122/20	3261/11	1268/20	2524/15
15	3063	0496/20	0618/4	0554/20	3045/4
16	118035	9168/19	3822/10	9305/19	5689/17
17	41004	8647/19	2601/9	8741/19	4356/15
18	23025	7855/19	5085/10	7884/19	2603/11
19	176035	7472/19	8046/6	758/19	4857/11
20	113044	4237/19	3814/11	442/19	7319/16
21	197017	2133/19	1349/11	2163/19	7429/14
22	181018	1163/19	4872/7	121/19	4562/10
23	35070	9636/18	6493/16	9703/18	8817/23
24	163014	7877/18	9464/10	7961/18	2093/16
25	101087	417/18	962/9	423/18	1953/11
26	157055	1672/18	8996/9	174/18	6058/11
27	242078	1324/18	006/11	1425/18	5446/16
28	42049	9737/17	8333/14	9908/17	2614/27
29	245051	3408/17	9813/13	3586/17	2115/26
30	35010	214/17	5393/12	2309/17	85/21

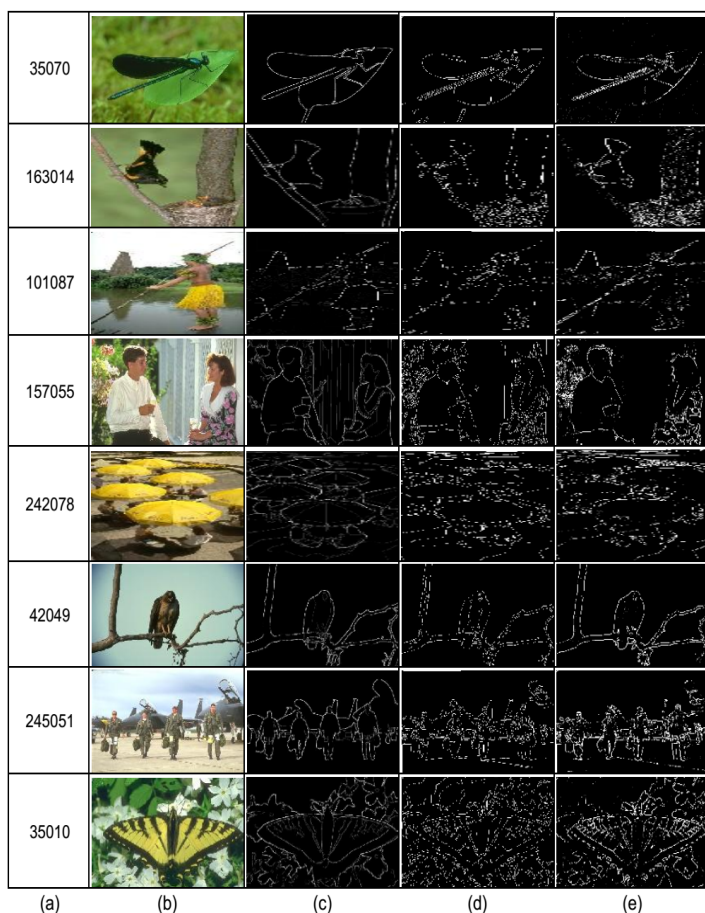


FIGURE 6: Column (a) Image No., Column (b) BSD image, Column (c) Ground Truth, Column (d) Canny's approach, Column (e) Proposed approach

شکل (۸): مجموعه تصاویر لبه‌یابی شده به روش پیشنهادی: ستون (a) شماره تصویر، ستون (b) تصویر RGB، ستون (c) Ground Truth، ستون (d) روش کنی، ستون (e) روش فازی پیشنهادی

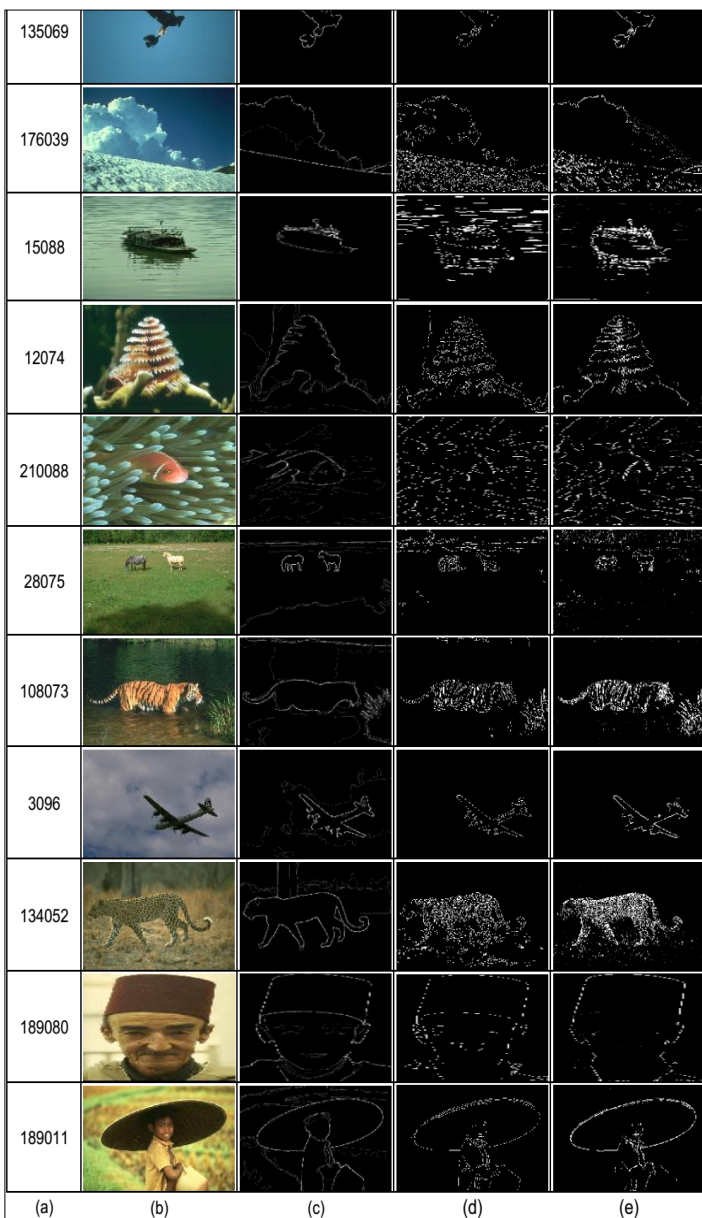


FIGURE 4: Column (a) Image No., Column (b) BSD image, Column (c) Ground Truth, Column (d) Canny's approach, Column (e) Proposed approach

۵- نتیجه گیری

تشخیص لبه یکی از تکنیک‌های مهم مورد استفاده برای تقسیم‌بندی تصویر است. تقسیم‌بندی تصویر حتی پس از چهار دهه همچنان یک مسئله مهم می‌باشد (منگ، ژانگ، ین، و ما، ۲۰۱۸). علاوه بر روش‌های سنتی در تشخیص لبه، در سال‌های اخیر روش‌ها و تکنیک‌های مختلفی به منظور لبه‌یابی ارائه شده است. عمده روش‌های ارائه شده بر اساس ترکیب روش‌های سنتی در لبه‌یابی است. روش پیشنهاد شده در این پژوهش به دلیل حفظ پیوستگی لبه و حذف نویز، نتایج بهتری نسبت به برخی از روش‌های معمول نشان می‌دهد. در این پژوهش با استفاده از مفهوم فازی روشی جدید برای لبه‌یابی ارائه شده است

که در آن تصویر به عنوان یک مجموعه فازی و پیکسل‌ها عناصر مجموعه فازی محسوب می‌شوند. روش فازی پیشنهادی تصویررنگی را به یک تصویرخوشه‌بندی شده تبدیل می‌کند، در نهایت یک تشخیص دهنده لبه بر روی تصویر خوشه‌بندی شده عمل کرده تا تصویر لبه‌دار حاصل شود. همان‌طور که عملگر پیشنهادی عمل تیره کردن^۲ را روی تصویر انجام نمی‌دهد، لبه‌های دوگانه را نیز کمتر شناسایی می‌کند. عموماً تصاویر واقعی لبه‌های قوی و ضعیف را شامل شده است. روش پیشنهادی هر دو لبه قوی و ضعیف را با قدرت لبه‌های مختلف با استفاده از حد آستانه بالا و پایین می‌دهد.

1- Meng, Zhang, Yin, & Ma,

2- Blurring

منابع:

- 1-Batan, M., Bukhari, S. S., & Breuel, T. (2017). Active Canny: edge detection and recovery with open active contour models. *IET Image Processing*, 11(12), 1325-1332.
- 2-Biswas, S., & Hazra, R. (2018). Robust edge detection based on Modified Moore-Neighbor. *Optik*, 168, 931-943.
- 3-Bozorgmehr, A., Jooq, M. K. Q., Moaiyeri, M. H., Navi, K., & Bagherzadeh, N. (2020). A novel digital fuzzy system for image edge detection based on wrap-gate carbon nanotube transistors. *Computers & Electrical Engineering*, 87, 106811.
- 4-Bozorgmehr, A., Moaiyeri, M. H., Navi, K., & Bagherzadeh, N. (2017). Ultra-efficient fuzzy min/max circuits based on carbon nanotube FETs. *IEEE Transactions on Fuzzy Systems*, 26(2), 1073-1078.
- 5-Dhivya, R., & Prakash, R. (2019). RETRACTED ARTICLE: Edge detection of satellite image using fuzzy logic. *Cluster Computing*, 22(Suppl 5), 11891-11898.
- 6-Driankov, D., Hellendoorn, H., & Reinfrank, M. (2013). *An introduction to fuzzy control*. Springer Science & Business Media.
- 7-Flores-Vidal, P. A., Olaso, P., Gómez, D., & Guada, C. (2019). A new edge detection method based on global evaluation using fuzzy clustering. *Soft Computing*, 23, 1809-1821.
- 8-Kalra, A., & Chhokar, R. L. (2016, September). A Hybrid approach using sobel and canny operator for digital image edge detection. In *2016 International Conference on Micro-Electronics and Telecommunication Engineering (ICMETE)* (pp. 305-310). IEEE.
- 9-Khaire, P. A., & Thakur, N. V. (2012). A fuzzy set approach for edge detection. *International Journal of Image Processing (IJIP)*, 6(6), 403-412.
- 10-Kumar, A., & Raheja, S. (2020). Edge detection using guided image filtering and enhanced ant colony optimization. *Procedia Computer Science*, 173, 8-17.
- 11-KumarSingh, R., Shekhar, S., Bhawan Singh, R., & Chauhan, V. (2014). A comparative study of edge detection techniques. *International Journal of Computer Applications*, 100(19), 5-8.
- 12-Meng, Y., Zhang, Z., Yin, H., & Ma, T. (2018). Automatic detection of particle size distribution by image analysis based on local adaptive canny edge detection and modified circular Hough transform. *Micron*, 106, 34-41.
- 13-Orujov, F., Maskelinis, R., Damaševičius, R., & Wei, W. J. A. S. C. (2020). Fuzzy based image edge detection algorithm for blood vessel detection in retinal images. *Applied Soft Computing*, 94, 106452.
- 14-Raheja, S., & Kumar, A. (2021). Edge detection based on type-1 fuzzy logic and guided smoothening. *Evolving Systems*, 12(2), 447-462.
- 15-Ren, T., Wang, H., Feng, H., Xu, C., Liu, G., & Ding, P. (2019). Study on the improved fuzzy clustering algorithm and its application in brain image segmentation. *Applied Soft Computing*, 81, 105503.
- 16-Qi, Y., Yang, Z., Sun, W., Lou, M., Lian, J., Zhao, W., ... & Ma, Y. (2021). A comprehensive overview of image enhancement techniques. *Archives of Computational Methods in Engineering*, 1-25.
- 17-Wang, A., Zhang, W., & Wei, X. (2019). A review on weed detection using ground-based machine vision and image processing techniques. *Computers and electronics in agriculture*, 158, 226-240.

©Authors, Published by Journal of Intelligent Knowledge Exploration and Processing. This is an open-access paper distributed under the CC BY (license <http://creativecommons.org/licenses/by/4.0/>).



مقاله پژوهشی

واکاوی تاثیر جهت‌گیری‌های کارآفرینانه سبز بر شیوه‌های مدیریت زنجیره تامین سبز

Doi: 10.30508/KDIP.2023.379963.1058

مسعود بختیاری آزاده (نویسنده مسئول)^۱ | بهاره شهریاری^۲^۱- گروه حسابداری، دانشگاه پیام نور (مرکز تهران)، تهران، ایران^۲- گروه مدیریت، دانشگاه پیام نور (مرکز تهران)، تهران، ایران

تاریخ دریافت: ۱۴۰۱/۱۰/۱۶

تاریخ پذیرش: ۱۴۰۱/۱۲/۰۶

صفحه: ۰۰ - ۰۰

چکیده:

جهت‌گیری استراتژیک یک حوزه مورد توجه و گسترده در ادبیات مدیریت کسب و کار است. فعالیت‌های استراتژیک شرکت می‌تواند معنای کلی جهت‌گیری استراتژیک را توصیف کند، که به طور مستقیم بر انجام عملکرد برتر شرکت تأثیر می‌گذارد. تحقیق حاضر؛ واکاوی تاثیر جهت‌گیری‌های کارآفرینانه سبز بر شیوه‌های مدیریت زنجیره تامین سبز در شرکت‌های شهرک صنعتی یزد می‌باشد. در این تحقیق جامعه آماری ۳۰ نفر از کارکنان شرکت‌های شهرک صنعتی یزد می‌باشد، که با استفاده از فرمول کوکران تعداد نمونه نهایی ۱۶۸ نفر می‌باشد. پس از توزیع پرسشنامه‌ها؛ تعداد ۱۵۰ پرسشنامه جمع‌آوری گردید. در این تحقیق از روش مدل‌سازی معادلات ساختاری با بکارگیری نرم‌افزار پی‌ا ال. اس. و اس. پی. اس. اس. استفاده شده است. نتایج نشان می‌دهد جهت‌گیری کارآفرینی سبز بر جهت‌گیری بازار، شیوه‌های مدیریت زنجیره تامین سبز و جهت‌گیری مدیریت دانش تأثیر معنی‌داری دارد. همچنین جهت‌گیری مدیریت دانش بر شیوه‌های مدیریت زنجیره تامین سبز تأثیر معنی‌داری دارد و جهت‌گیری کارآفرینی سبز بر شیوه‌های مدیریت زنجیره تامین سبز از طریق جهت‌گیری مدیریت دانش تأثیر دارد. فعالیت‌های استراتژیک شرکت، می‌تواند معنای کلی جهت‌گیری استراتژیک را توصیف کند، که به طور مستقیم بر انجام عملکرد برتر شرکت تأثیر می‌گذارد.

کلمات کلیدی: جهت‌گیری‌های استراتژیک، مدیریت زنجیره تامین، زنجیره تامین سبز، عملکرد پایدار، کارآفرینی سبز.

۱- مقدمه

تحقیقات متعددی ارتباط بین جهت‌گیری استراتژیک و عملکرد شرکت را بررسی کرده‌اند (گست، گاندولف، و کسینگر، ۲۰۱۷). البته شواهد خاصی در مورد اینکه چگونه جهت‌گیری استراتژیک عملکرد را به ارمغان می‌آورد، وجود ندارد (جیانگ، چای، شائو و فنگ، ۲۰۱۸). به عنوان مثال، اخیراً گست و همکاران (۲۰۱۷) به شناسایی فعالیت‌های جهت‌گیری استراتژیک که بر عملکرد اقتصادی و زیست محیطی شرکت تأثیر مثبت دارد، پرداختند. بر اساس بررسی‌های به عمل آمده، مطالعات محدودی وجود دارد که به طور سیستماتیک راه‌های استراتژیک دستیابی به عملکرد برتر شرکت را با در نظر گرفتن اطمینان از پایداری، به ویژه از نظر کسب و کار و تولید بررسی کرده باشد. رشد پایدار صنایع و در مقابل تأثیر منفی قابل توجه بر پایداری محیط زیست، یک چالش بزرگ است و تولیدکنندگان در تلاش برای اجرای شیوه‌های مدیریت زنجیره تامین سبز و بهبود عملکرد شرکت هستند (اسلام، پری و گیل، ۲۰۲۱). بنابراین، مطالعات کمی در خصوص بررسی ارتباط بین جهت‌گیری استراتژیک و عملکرد شرکت وجود دارد (هاگست، بائو، و ایمودین، ۲۰۱۸). از این رو، مطالعه ما عوامل تأثیرگذار جهت‌گیری استراتژیک را بررسی می‌کند، مانند اتخاذ شیوه‌های مدیریت زنجیره تامین سبز، که عملکرد پایدار شرکت برتر را به ارمغان می‌آورد (حبیب، هادکینسون، هاگز و ارشد، ۲۰۲۰). همچنین، این مطالعه بررسی می‌کند، چگونه جهت‌گیری‌های بازار و مدیریت دانش نقش واسطه‌ای در تأثیر جهت‌گیری کارآفرینی سبز بر مدیریت زنجیره تامین سبز دارند.

استان یزد یکی از صنعتی‌ترین مناطق ایران بوده و بسیاری از روزهای سال، هوای شهر آلودگی بالایی را تجربه می‌کند. با این حال، اجرای مدیریت زنجیره تامین سبز بدون انگیزه، فشار، پشتیبانی مدیریت ارشد، زیرساخت سبز، عرضه سبز و تسهیلات سیستم اطلاعات سبز

آسان نیست. فشار دولت، فرهنگ سازمانی، دانش سبز، توانایی فن‌آوری و مالی، مزایای اقتصادی و اجتماعی حاصل از شیوه‌های سبز ممکن است زمینه‌ای برای پذیرش مدیریت زنجیره تامین سبز در صنعت سیم و کابل باشد (ماجومدار، و سینها، ۲۰۱۸). اگرچه برخی از شرکت‌ها گواهینامه‌های ایزو ۹۰۰۰^۷ و ایزو ۱۴۰۰۱^۸ را پذیرفته‌اند، اما همچنان، بسیاری از صنایع از عملکرد مدیریت زنجیره تامین سبز عقب مانده‌اند. از این رو، سازمان‌ها باید در اتخاذ شیوه‌های مدیریت زنجیره تامین سبز برای بهبود عملکرد پایدار استراتژیک باشند.

جهت‌گیری استراتژیک از اصول هدایت‌کننده تأثیرگذار بر بازاریابی و انتخاب استراتژی‌های فعالیت‌های یک شرکت هستند که منعکس‌کننده گرایش‌های استراتژیک اجرا شده توسط شرکت برای ایجاد رفتارهای مناسب می‌باشند، به عملکرد بهتر منجر می‌شوند، و بر اساس اندیشه شرکت در مورد انجام کسب و کار از طریق مجموعه‌ای گسترده از ارزش‌ها و باورهای ریشه‌ای شکل گرفته‌اند. در تجارت جهانی امروز، رقابت میان سازمان‌ها بسیار شدید است و برای تحت تأثیر قرار دادن مشتریان، سازمان‌ها نیاز دارند خودشان را در موقعیت برتری نسبت به رقبای خود قرار دهند. دواستدار محیط زیست بودن و سازگاری با الزامات زیست محیطی، راهی برای تمایز از سایر رقبا است. در صورتی که رقبا از مدیریت زنجیره تامین سبز بهره‌مند شده باشند، شرکت تحت فشار بیشتری برای استقرار مدیریت زنجیره تامین سبز خواهد بود. از طرفی مشتریان نیز روی تصمیم برای استقرار سیستم مدیریت زنجیره تامین سبز نقش مهمی دارند. از این رو، مدیریت استراتژیک زنجیره تامین سبز، یکی از الزامات موفقیت در بازار رقابتی امروز است و باید به طور جدی مدنظر قرار گیرد (احمدی نژاد، کریمی زارچی، و فتاحی، ۱۳۹۹).

با توجه به اینکه مطالعات قبلی رابطه بین جهت‌گیری استراتژیک و عملکرد تجاری را بررسی کرده‌اند (فیصل،

1- Gast, Gundolf, & Cesinger

2- Jiang, Chai, Shao, & Feng

3- Islam, M. M., Perry, P., & Gill

4- Hughes, Hodgkinson, Hughes, & Arshad

5- Habib, Bao, & Ilmudeen

6- Majumdar, & Sinha

7- ISO9000

8- ISO14001

هرماوان، و ارافاه^۱، از سوی دیگر، بسیاری از مطالعات نشان دادند که اجرای شیوه‌های مدیریت زنجیره تامین سبز به عملکرد پایداری شرکت کمک می‌کند (الطیب، زایلانی، ورمایه^۲، ۲۰۱۱). با این حال، تحقیقات کمی در خصوص ارزیابی جهت‌گیری استراتژیک و شیوه‌های مدیریت زنجیره تامین سبز در مورد عملکرد اقتصادی و محیطی سازمانی وجود دارد. لذا، دیدگاه کلی جهت‌گیری استراتژیک در اجرای بهترین شیوه‌های مدیریت زنجیره تامین سبز در مورد عملکرد پایداری فاقد ادبیات نظری و عملی است. علاوه بر این، تحقیقات قبلی در حوزه مدیریت زنجیره تامین سبز عمدتاً به تجزیه و تحلیل اثر مستقیم ساده (فوو، لی، تان، اووی^۳، ۲۰۱۸)، یا شناسایی مؤلفه‌ها و اهمیت آنها (لوترا، گارج، حلیم^۴، ۲۰۱۵) بوده است. این مطالعه اولین تلاش برای شناسایی تأثیرات مدیریت زنجیره تامین سبز در قالب جهت‌گیری استراتژیک می‌باشد. بر این اساس مقاله حاضر، واکاوی تأثیر جهت‌گیری‌های کارآفرینانه سبز بر شیوه‌های مدیریت زنجیره تامین سبز در شرکت‌های شهرک صنعتی یزد می‌باشد، که اهداف فرعی زیر را نیز دنبال داشت:

- ۱- شناسایی تأثیر جهت‌گیری کارآفرینی سبز بر جهت‌گیری بازار
- ۲- شناسایی تأثیر جهت‌گیری کارآفرینی سبز بر شیوه‌های مدیریت زنجیره تامین سبز
- ۳- شناسایی تأثیر جهت‌گیری کارآفرینی سبز بر جهت‌گیری مدیریت دانش
- ۴- شناسایی تأثیر جهت‌گیری بازار بر شیوه‌های مدیریت زنجیره تامین سبز
- ۵- شناسایی تأثیر جهت‌گیری مدیریت دانش بر شیوه‌های مدیریت زنجیره تامین سبز
- ۶- شناسایی تأثیر جهت‌گیری کارآفرینی سبز بر شیوه‌های مدیریت زنجیره تامین سبز از طریق جهت‌گیری بازار

۷- شناسایی تأثیر جهت‌گیری کارآفرینی سبز بر شیوه‌های مدیریت زنجیره تامین سبز از طریق جهت‌گیری مدیریت دانش

با توجه به اثرات مخرب زیست محیطی و اجتماعی در عملیات زنجیره تامین در کارخانجات و صنایع تولیدی، تغییرات اساسی در مدل‌های کسب و کار در جهت حرکت به سمت توسعه پایدار از اهمیت زیادی برخوردار است (کوبرگ، ولانگیون^۵، ۲۰۱۹). با توجه به افزایش آگاهی زیست محیطی جهانی، امروزه شیوه‌های مدیریت زنجیره تامین سبز ضروری و یکی از موثرترین شیوه‌های پایداری در میان کلیه سازمان‌ها و صنایع تولیدی و خدماتی؛ دولتی و غیردولتی شد (لین، لوو، لرومانچو، رانگ، هانگ^۶، ۲۰۱۹). تحقیقات اشاره می‌کند، مدیریت صنایع تولیدی با توجه به ماهیت پیچیده عملیات، در حال تلاش برای رسیدگی به مسائل پایداری در زنجیره تامین کسب و کار خود هستند (نینیمکی، و همکاران^۷، ۲۰۲۰). جهت‌گیری استراتژیک یک حوزه مورد توجه و گسترده‌تر در ادبیات مدیریت کسب و کار است. فعالیت‌های استراتژیک شرکت می‌تواند معنای کلی جهت‌گیری استراتژیک را توصیف کند، که به طور مستقیم بر انجام عملکرد بزرگ شرکت تأثیر می‌گذارد (ساهی، گوپتا، چنگ^۸، ۲۰۱۹). در طول این سال‌ها، برخی از محققان انواع مختلف جهت‌گیری استراتژیک را شناسایی کرده‌اند که پیامدهای چند بعدی آن‌ها را در سطح سازمان نشان می‌دهد. در این مطالعه جهت‌گیری استراتژیک شامل سه بعد؛ کارآفرینی سبز، بازار و مدیریت دانش است. این سه بعد، بالاترین قابلیت تأثیرگذاری را برای اتخاذ فعالیت‌های محیطی سازمانی داشته و به عملکرد بزرگ شرکت کمک می‌کند (ایورس، گلیگا، وریالپ کریادو^۹، ۲۰۱۹).

1- Faisal, Hermawan, & Arafah

2- Eltayeb, Zailani, & Ramayah

3- Foo, Lee, Tan, & Ooi

4- Luthra, Garg, & Haleem

5- Koberg, & Longoni

6- Lin, Luo, Jeromonachou, Rong, & Huang

7- Niinimäki et al

8- Sahi, Gupta, & Cheng

9- Evers, Gliga, & Rialp-Criado

۲- مبانی نظری

جهت‌گیری کارآفرینی سبز

جهت‌گیری کارآفرینی سبز به عنوان یک چشم‌انداز برای تغییر جامعه از طریق کاهش پیشگیری از آلودگی و کاهش اثرات زیست محیطی تعریف می‌شود. جهت‌گیری کارآفرینی سبز، بدون شک ترکیبی از برخی ویژگی‌های کارآفرینی به عنوان یک وضعیت تصمیم‌گیری در جهت فرآیند تصمیم‌گیری است (گست و همکاران، ۲۰۱۷). کارآفرینان سبز نیز با هدف نوآوری سبز، ریسک‌پذیری و جستجوی فرصت‌های جدید، سود می‌برند و فرصت‌های جدیدی را پیدا می‌کنند (پاچئو، دین و پاینه، ۲۰۱۰). به عقیده دین و مک مولن^۲، جهت‌گیری کارآفرینی سبز تمایل به شناسایی فرصت‌های بالقوه‌ای است که اغلب با شروع فعالیت‌های سبز، رفاه اقتصادی و زیست محیطی ایجاد می‌شود. معمولاً جهت‌گیری کارآفرینی، موضع تصمیم‌گیری شرکت را در وظایف حیاتی در سطح شرکت، فرآیند تصمیم‌گیری استراتژی و ایده‌های مدیریتی برای یافتن فرصت‌های جدید برای رشد و تجدید سازمان درک می‌شود (دین و مک مولن^۳، ۲۰۰۷).

جهت‌گیری بازار

جهت‌گیری بازار به عنوان استراتژی‌ها و عملیاتی از شرکت تعریف می‌شود که از آن برای برآورده کردن کاربرد فعلی و آتی محصول و بازار بهره گرفته می‌شود (گرین، زلبست، میچام، و بهادوریا^۴، ۲۰۱۲). جهت‌گیری بازار را می‌توان به عنوان توانایی شناسایی و ارضای نیازهای مشتریان تعریف کرد. این مطالعه دو دیدگاه رفتاری جهت‌گیری بازار را در قالب: مشتری‌گرایی و رقیب‌گرایی در نظر می‌گیرد. اخیراً، با رشد نگرانی‌های زیست محیطی، مشتریان خواستار محصولات دوست‌دار محیط زیست هستند (فرامباچ، پرابو، و ورهالان^۵، ۲۰۰۳). شرکت‌های بازارگرا چنین نیازهایی را جلب نموده و

منابع شرکت را به سمت طرح‌های سبز بسیج می‌کنند. از آنجایی که شرکت بازارگرا از دانش بازار بالاتری نسبت به رقبای خود برخوردار است. شرکت‌های بازارمحور همیشه به طور سیستماتیک جمع‌آوری می‌شوند، استراتژی رقیب را پایش، تحلیل می‌کنند و برای دستیابی به مزیت‌های رقابتی اقدام سبز انجام می‌دهند. جهت‌گیری بازار مزایای بازاریابی را از طریق وابستگی پایدار شیوه‌های مدیریت زنجیره تامین سبز افزایش می‌دهد. جهت‌گیری بازار یک منبع نامشهود است که دانش مدیریت را در مورد تقاضای مشتری برای محصول زیست محیطی بهبود می‌بخشد (چوی^۶، ۲۰۱۴).

جهت‌گیری مدیریت دانش

دانش یک منبع استراتژیک سازمانی است که شرکت‌ها را برای دستیابی به عملکرد رقابتی راهنمایی می‌کند. از این رو، سازمان باید بر توسعه درک خود از حمایت مبتنی بر جمع‌آوری دانش بازار تأکید کند. یادگیری یک منبع استراتژیک است و برای کسب دانش بهتر از مشتری و رقیب برای ایجاد موقعیت رقابتی در بازار استفاده می‌کند (زیام، کسلی، جارادات، و کسترات^۷، ۲۰۱۸). تجربه منبعی ارزشمند، کمیاب، تکرار نشدنی و غیرقابل جایگزینی است که آنها را به یک شرکت فناپذیرتر تبدیل می‌کند. در فرآیند زنجیره تامین، آنها باید دانش و اطلاعات را بین شرکای زنجیره تامین به اشتراک بگذارند تا عملکردهای مختلف را به پایان برسانند و تقاضای مشتری را برآورده کنند. به همین خاطر؛ جهت‌گیری مدیریت دانش اطلاعات تامین کنندگان و مشتریان سازمان را فراهم می‌کند. جهت‌گیری مدیریت دانش به سازمان کمک می‌کند تا ارتباط مناسب با اطلاعات به روز را حفظ کند و همچنین روابط خوب با تامین کنندگان خود را برای بهبود رشد کسب و کار توسعه و مدیریت کند (تسنگ^۸، ۲۰۱۴).

1- Pacheco, Dean, & Payne

2- Dean & McMullen

3- Dean, & McMullen

4- Green, Zelbst, Meacham, & Bhadauria

5- Frambach, Prabhu, & Verhallen

6- Choi

7- Zaim, Keceli, Jaradat, & Kastrat

8- Tseng

شیوه‌های مدیریت زنجیره تامین سبز

مدیریت زنجیره تامین هماهنگی در تولید، موجودی، مکان‌یابی و حمل و نقل بین شرکت کنندگان در یک زنجیره تامین است که جهت دستیابی به بهترین ترکیب پاسخگویی و کارایی برای موفقیت در بازار بکار می‌رود (کاپاسینو^۱، ۱۹۹۷). مدیریت زنجیره تامین سبز از منظر چرخه عمر محصول شامل تمامی مراحل از مواد اولیه، طراحی و ساخت محصول، فروش محصول و حمل و نقل، استفاده از محصول و بازیافت محصولات می‌باشد. با استفاده از مدیریت زنجیره تامین و فناوری سبز، شرکت می‌تواند تأثیرات منفی زیست محیطی را کاهش داده و به استفاده مطلوب از منابع و انرژی دست یابد (تسنگ، اسلام، کاریا، فوزی و افرین^۲، ۲۰۱۹). سبز کردن زنجیره تامین، فرایند در نظر گرفتن معیارها یا ملاحظات زیست محیطی در سرتاسر زنجیره تامین است. مدیریت زنجیره تامین سبز، یکپارچه کننده مدیریت زنجیره تامین با الزامات زیست محیطی در تمام مراحل طراحی محصول، انتخاب و تامین سبز کردن زنجیره تامین، فرایند در نظر گرفتن معیارها یا

ملاحظات زیست محیطی در سرتاسر زنجیره تامین است. لذا؛ شیوه‌های مدیریت زنجیره تامین سبز نیز به شیوه هایی گفته می‌شود که در آن مفاهیم پایداری را در مدیریت زنجیره تامین سنتی گنجانده است. در این شیوه، هدف کمک به صنایع برای کاهش انتشار کربن و به حداقل رساندن ضایعات و در عین حال به حداکثر رساندن سود است (گرین و همکاران، ۲۰۱۲). در ادامه به برخی از تحقیقات پیشین، مدل مفهومی تحقیق و سپس، فرضیات تحقیق اشاره شده است.

جهت‌گیری استراتژیک از اصول هدایت کننده تأثیرگذار بر بازاریابی و انتخاب استراتژی‌های فعالیت‌های یک شرکت هستند که منعکس کننده گرایش‌های استراتژیک اجرا شده توسط شرکت برای ایجاد رفتارهای مناسب می‌باشند و به عملکرد بهتر منجر می‌شوند. بر اساس اندیشه شرکت در مورد انجام کسب و کار از طریق مجموعه‌ای گسترده از ارزش‌ها و باورهای ریشه‌ای شکل گرفته‌اند. در جدول شماره (۱) پیشینه تحقیق بیان شده است.

جدول (۱): پیشینه تحقیق		
محققین	عنوان	نتایج
کیانی، مروتی شریف آبادی، مفتاح زاده، و زمزم (۱۴۰۱)	طراحی مدل پویای عوامل موثر بر موفقیت مدیریت زنجیره تامین سبز	تقویت مهارت نیروی انسانی، نظارت بر اجرای قوانین و مقررات و تامین کنندگان سبز، به طور مستقیم منجر به افزایش موفقیت مدیریت زنجیره تامین خواهد شد؛ که در این بین نظارت بر اجرای قوانین و مقررات مؤثرترین عامل شناخته شده است. متغیر سفارش برای تولید سبز برای بقای زنجیره تامین سبز بسیار حائز اهمیت است. این مدل و شبیه سازی آن می‌تواند به عنوان یک سیستم پشتیبانی تصمیم کاربردی برای مدیران سازمان‌ها باشد تا با توجه به شرایط محیطی خود بهترین تصمیمات را بگیرند.

1- Copacino

2- Tseng, Islam, Karia, Fauzi, & Afrin

جدول (۱): پيشينه تحقيق		
محققين	عنوان	نتايج
عظيمي، موسوي، و خورشيدى (۱۴۰۰)	بررسى تاثير بازارگرایی بر عملکرد محيط زيستى با نقش ميانجى راهكارهاى مديريت زنجيره تامين سبز	بازارگرایی بر راهكارهاى مديريت زنجيره تامين موثر بود. اما تاثير آن بر عملکرد محيط زيستى رد شد. از سوى ديگر، تاثير غيرمستقيم بازارگرایی از طريق متغير ميانجى راهكارهاى مديريت زنجيره تامين تايد شد. همچنين راهكارهاى مديريت زنجيره تامين بر عملکرد محيط زيستى تاثير مثبت و معنادارى دارد.
ايرجى راد و پارسامهر (۱۴۰۰)	تاثير جهت‌گيرى استراتژيک حلقه بسته بر عملکرد با تبیین نقش ميانجى مديريت زنجيره تامين سبز در شرکت ايران خودرو	جهت‌گيرى استراتژيک حلقه بسته و اقدامات مديريت زنجيره تامين سبز، بر عملکرد زيست محيطى، عملکرد اقتصادى مثبت، عملکرد اقتصادى منفى و متغير اقدامات مديريت زنجيره تامين سبز تاثير مثبت و معنى دارى دارد. اقدامات مديريت زنجيره تامين سبز، نقش ميانجى در رابطه بين جهت‌گيرى استراتژيک حلقه بسته با عملکرد زيست محيطى، عملکرد اقتصادى مثبت، عملکرد اقتصادى منفى دارد.
حبيب و همکاران (۲۰۲۱)	تاثير جهت‌گيرى‌هاى استراتژيک بر اجراى شيوه‌هاى مديريت زنجيره تامين سبز و عملکرد پايدار شرکت	جهت‌گيرى‌هاى مديريت دانش تاثير مثبتى بر شيوه‌هاى شيوه‌هاى مديريت زنجيره تامين سبز ندارد. و تا حدى رابطه بين جهت‌گيرى‌هاى کارآفرينانه سبز و عملکرد شرکت پايدار را واسطه‌گرى مى‌کنند. در حالى که جهت‌گيرى‌هاى بازار و مديريت دانش تا حدى رابطه بين عملکردهاى جهت‌گيرى‌هاى کارآفرينانه سبز را واسطه مى‌کنند.
برازون، هانگ و ليو (۲۰۲۱)	جهت‌گيرى بازار سبز و عملکرد سازمانى در صنعت برق و الکترونیک تايوان: نقش واسطه‌اى قابليت مديريت زنجيره تامين سبز	جهت‌گيرى بازار سبز تاثير مثبت معنى دارى بر قابليت مديريت زنجيره تامين سبز، عملکرد زيست محيطى و عملکرد اقتصادى دارند. قابليت مديريت زنجيره تامين سبز به طور مثبت با عملکرد زيست محيطى و اقتصادى مرتبط است. همچنين جهت‌گيرى بازار سبز از طريق قابليت مديريت زنجيره تامين سبز تاثير غيرمستقيم قابل توجهى بر عملکرد زيست محيطى و عملکرد اقتصادى دارند.
گروزا، زمج، و رجبى (۲۰۲۱)	نوآوری سازمانى و عملکرد شرکت: آیا مديريت فروش اهميت دارد؟	انگيزش ذهنى مدير فروش منجر به نوآوری سازمانى مى‌شود. رابطه بين نوآوری سازمانى و رشد فروش از شکل U غيرخطى پيروي مى‌کند.

فرضيه اصلى تحقيق

جهت‌گيرى‌هاى استراتژيک بر اجراى شيوه‌هاى مديريت زنجيره تامين سبز و عملکرد پايدار در شرکت‌هاى شهرک صنعتى يزد تاثير دارد.

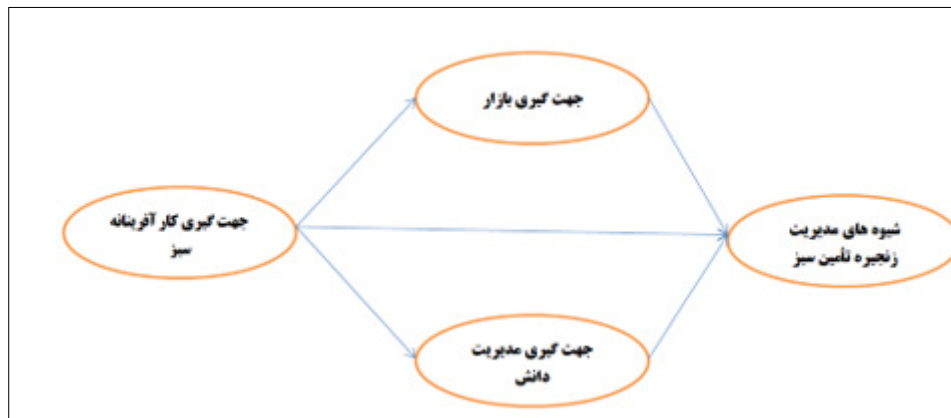
فرضيه‌هاى فرعى تحقيق

- ۸- جهت‌گيرى کارآفرينى سبز بر جهت‌گيرى بازار تاثير معنى دارى دارد.
- ۹- جهت‌گيرى کارآفرينى سبز بر شيوه‌هاى مديريت زنجيره تامين سبز تاثير معنى دارى دارد.
- ۱۰- جهت‌گيرى کارآفرينى سبز بر جهت‌گيرى مديريت

دانش تاثير معنى دارى دارد.

- ۱۱- جهت‌گيرى بازار بر شيوه‌هاى مديريت زنجيره تامين سبز تاثير معنى دارى دارد.
 - ۱۲- جهت‌گيرى مديريت دانش بر شيوه‌هاى مديريت زنجيره تامين سبز تاثير معنى دارى دارد.
 - ۱۳- جهت‌گيرى کارآفرينى سبز بر شيوه‌هاى مديريت زنجيره تامين سبز از طريق جهت‌گيرى بازار تاثير دارد.
 - ۱۴- جهت‌گيرى کارآفرينى سبز بر شيوه‌هاى مديريت زنجيره تامين سبز از طريق جهت‌گيرى مديريت دانش تاثير دارد.
- با توجه به مبانى نظرى و فرضيه‌هاى تحقيق، مى‌توان مدل

مفهومی تحقیق را در قالب شکل شماره (۱) به شرح زیر ارائه نمود.



شکل (۱): مدل مفهومی پژوهش (حبیب و همکاران، ۲۰۲۱)

۳- روش تحقیق

اندازه‌گیری شد. شیوه‌های مدیریت زنجیره تامین سبز از طریق نه سوال بر اساس مطالعه ژو، سارکیس، و لای^۳ (۲۰۰۷)؛ عملکرد اقتصادی و زیست محیطی نیز هرکدام از طریق ۵ سوال بر اساس مطالعه پائولراج^۴ (۲۰۱۱) اندازه‌گیری گردید. در تجزیه و تحلیل داده‌های این تحقیق، از شاخص‌های آمار توصیفی شامل؛ فراوانی، درصد، میانگین و انحراف معیار برای توصیف متغیرها و از آزمون‌های تحلیلی جهت بررسی فرضیه‌ها استفاده شد. همچنین؛ از روش مناسب مدل‌سازی معادلات ساختاری (SEM) و نرم‌افزار پی.ا.ال.اس. و اس.پی.اس.اس. استفاده شده است. جهت پایایی و روایی پرسشنامه، جدول‌های شماره (۲)، (۳) و (۴) در ادامه ارائه گردیده است.

مطالعه حاضر از نوع تحلیلی و برحسب زمانی از نوع مقطعی می‌باشد که به صورت پیمایشی بر روی کارکنان شرکت‌های شهرک صنعتی یزد به انجام رسید. در این تحقیق جامعه آماری؛ ۳۰۰ نفر از کارکنان شرکت‌های شهرک صنعتی یزد می‌باشد که با استفاده از فرمول کوکران تعداد نمونه نهایی ۱۶۸ نفر تعیین گردید. پس از توزیع پرسشنامه‌ها، تعداد ۱۵۰ پرسشنامه جمع‌آوری شد. روش نمونه‌گیری به صورت تصادفی ساده بود. برای جهت‌گیری کارآفرینی سبز، ۵ سوال از مطالعه جیانگ و همکاران (۲۰۱۸)، جهت‌گیری بازار از طریق ۵ سوال برگرفته از مطالعه فرامباچ^۱ و همکاران (۲۰۰۳) و جهت‌گیری مدیریت دانش از طریق ۵ سوال برگرفته از مطالعه آتیا، و اسام الدین^۲ (۲۰۱۸)

جدول (۲): مقادیر ضریب آلفای کرونباخ و پایایی ترکیبی			
ردیف	سازه	آلفای کرونباخ	پایایی ترکیبی
۱	شیوه‌های مدیریت زنجیره تامین سبز	۰/۸۸۵	۰/۹۰۸
۲	جهت‌گیری کارآفرینی سبز	۰/۸۳۷	۰/۸۸۵
۳	جهت‌گیری مدیریت دانش	۰/۷۲۵	۰/۸۲۲
۴	جهت‌گیری بازار	۰/۸۱۱	۰/۸۶۹

1- Frambach

2- Attia, & Essam Eldin

3- Zhu, Sarkis, & Lai

4- Paulraj

در جدول شماره (۲)، ضريب آلفای کرونباخ و پايایى ترکیبى پايایى مناسب مدل دارد. برای کلیه متغیرهای مورد نظر بالاتر از ۰/۷ است که حاکی از

جدول (۳): مقادير روايی همگرا (AVE) ابعاد تحقيق		
ردیف	سازه	روايی همگرا (AVE)
۱	شیوه‌های مدیریت زنجیره تامین سبز	۰/۵۲۵
۲	جهت‌گیری کارآفرینی سبز	۰/۶۰۵
۳	جهت‌گیری مدیریت دانش	۰/۴۹۵
۴	جهت‌گیری بازار	۰/۵۷۲

لازم به ذکر است با توجه به نتایج داده‌های جدول شماره (۳)، در کلیه متغیرهای تحقیق (بجز جهت‌گیری مدیریت دانش)، مقادير روايی همگرایى (AVE) سازه‌های تحقیق، بیشتر از ۰/۵ می‌باشد، که بیان‌گر روايی مناسب متغیرهای تحقیق می‌باشد. در این تحقیق جهت بررسی روايی و اگر از روش فورنل و لارکر استفاده شده است.

جدول (۴): روايی و اگر روش فورنل و لارکر				
جهت‌گیری بازار	جهت‌گیری مدیریت دانش	جهت‌گیری کارآفرینی سبز	شیوه‌های مدیریت زنجیره تامین سبز	
			۰/۷۲۴	شیوه‌های مدیریت زنجیره تامین سبز
		۰/۷۷۸	۰/۵۹۴	جهت‌گیری کارآفرینی سبز
	۰/۶۰۴	۰/۶۰۹	۰/۷۳۰	جهت‌گیری مدیریت دانش
۰/۷۵۷	۰/۶۱۵	۰/۶۹۲	۰/۵۵۵	جهت‌گیری بازار

۴- یافته‌های تحقیق

از شاخص‌های آمار توصیفی برای بررسی ویژگی‌های دموگرافیک پاسخ‌دهندگان استفاده شده است. فراوانی پاسخ‌دهندگان بر اساس جنسیت، تاهل، سن، تحصیلات و سابقه کار بررسی قرار گرفته است که در جدول‌های شماره (۵) تا (۹) آورده شده است.

مطابق داده‌های جدول شماره (۴)، مجذور روايی همگرای هر سازه از مقادير همبستگی بین سازه‌های دیگر بزرگ‌تر می‌باشد. لذا؛ مدل تحقیق، از نظر روايی و اگر مطابق روش فورنل و لارکر مورد تأیید می‌باشد.

جدول ۵: فراوانی و درصد افراد مورد مطالعه در متغير جنسيت

جنسيت	فراوانی	درصد	درصد تجمعی
مرد	۱۱۱	۷۴	۷۴
زن	۳۹	۲۶	۱۰۰
کل	۱۵۰	۱۰۰	

جدول ۶: فراوانی و درصد افراد مورد مطالعه در متغير وضعیت تاهل

وضعيت تاهل	فراوانی	درصد	درصد تجمعی
مجرد	۲۸	۱۸٫۷	۱۸٫۷
متاهل	۱۲۲	۸۱٫۳	۱۰۰
کل	۱۵۰	۱۰۰	

جدول ۷: فراوانی و درصد افراد مورد مطالعه در متغير سن

سن	فراوانی	درصد	درصد تجمعی
کمتر از ۳۰ سال	۳۴	۲۲٫۷	۲۲٫۷
بين ۳۱ تا ۴۰ سال	۶۴	۴۲٫۷	۶۵٫۳
بين ۴۱ تا ۵۰ سال	۳۶	۲۴	۸۹٫۳
بالای ۵۰ سال	۱۶	۱۰٫۶	۱۰۰
جمع کل	۱۵۰	۱۰۰	

جدول ۸: فراوانی و درصد افراد مورد مطالعه در متغير تحصیلات

تحصیلات پاسخگویان	فراوانی	درصد	درصد تجمعی
ديپلم	۳۰	۲۰	۲۰
فوق ديپلم	۴۱	۲۷٫۳	۴۷٫۳
لیسانس	۵۱	۳۴	۸۱٫۳
فوق لیسانس	۲۲	۱۴٫۷	۹۶
دکتری	۶	۴	۱۰۰
جمع کل	۱۵۰	۱۰۰	

جدول ۹: فراوانی و درصد افراد مورد مطالعه در متغیر سابقه کاری			
سابقه کار	فراوانی	درصد	درصد تجمعی
کمتر از ۱۰ سال	۶۴	۴۲٫۷	۴۲٫۷
بین ۱۱ تا ۲۰ سال	۴۷	۳۱٫۳	۷۴
بیشتر از ۲۰ سال	۳۹	۲۶	۱۰۰
جمع کل	۱۵۰	۱۰۰	

جهت تحلیل توصیفی متغیرهای تحقیق از میانگین و انحراف معیار مطابق جدول شماره (۱۰) استفاده شده است.

جدول ۱۰: شاخص‌های توصیفی برای هریک از متغیرهای مطالعه					
متغیر	فراوانی	کمترین	بیشترین	میانگین نمرات	انحراف معیار
جهت‌گیری کارآفرینی سبز	۱۵۰	۱٫۶	۵	۳٫۴۳	۰٫۷۴
جهت‌گیری بازار	۱۵۰	۱٫۶	۵	۳٫۶۰	۰٫۷۴
جهت‌گیری مدیریت دانش	۱۵۰	۱٫۴	۵	۳٫۵۳	۰٫۶۴
شیوه‌های مدیریت زنجیره تامین سبز	۱۵۰	۲	۵	۳٫۶۲	۰٫۶۵

جهت تعیین نرمالیده بودن توزیع داده‌ها از آزمون‌های کولموگروف-اسمیرنوو و آزمون شاپرو-ویلک استفاده شد. جدول شماره (۱۱) نتایج این آزمون‌ها را نشان می‌دهد. از آنجایی که سطح معنی‌داری در کلیه متغیرها در حداقل یکی از دو آزمون کمتر از ۰٫۰۵ است، لذا می‌توان گفت؛ توزیع داده‌ها در متغیرهای مورد بررسی نرمال نیست.

جدول ۱۱: نتایج آزمون نرمالیده بودن داده‌های پژوهش			
متغیرها	فراوانی	معنی‌داری کولموگروف - اسمیرنوف	معنی‌داری شاپرو - ویلک
جهت‌گیری کارآفرینی سبز	۱۵۰	۰٫۰۰۰	۰٫۰۰۱
جهت‌گیری بازار	۱۵۰	۰٫۰۰۰	۰٫۰۰۱
جهت‌گیری مدیریت دانش	۱۵۰	۰٫۰۰۰	۰٫۰۰۰
شیوه‌های مدیریت زنجیره تامین سبز	۱۵۰	۰٫۰۳۶	۰٫۱۹۹

کیفیت مدل‌های اندازه‌گیری در روش پی.ال.اس.، با استفاده از معیار Commuality ارزیابی می‌شود. این معیار نشان می‌دهد که چه مقدار از تغییرپذیری شاخص‌ها (سوالات) توسط سازه‌ی مرتبط با خود تبیین می‌شود. مقادیر اشتراکی متعلق به سازه‌های مدل تحقیق در جدول شماره (۱۲) نشان داده شده که نشان دهنده کیفیت مدل اندازه‌گیری شده در این تحقیق می‌باشد.

جدول (۱۲): مقادیر اشتراکی ابعاد تحقیق		
ردیف	سازه	مقادیر اشتراکی
۱	شیوه‌های مدیریت زنجیره تامین سبز	۰/۳۹۳
۲	جهت‌گیری کارآفرینی سبز	۰/۴۰۰
۳	جهت‌گیری مدیریت دانش	۰/۲۷۱
۴	جهت‌گیری بازار	۰/۳۵۶

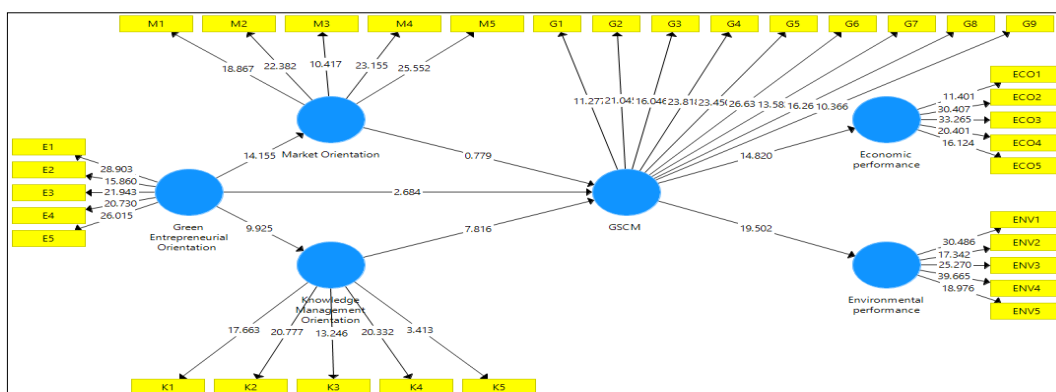
برای بررسی برازش مدل ساختاری تحقیق از ضرایب معنی‌داری t استفاده شد که در جدول شماره (۱۳) نمایش داده شده است. بر اساس نتایج این جدول مقادیر t برای شش فرضیه مستقیم بالاتر از ۱/۹۶ می‌باشد.

جدول ۱۳: مقادیر t -value برای فرضیه‌های مستقیم				
سطح معنی‌داری	مقادیر t	انحراف معیار	بار عاملی	
0.008	2.684	0.076	0.205	Green Entrepreneurial Orientation → GSCM
0.000	9.925	0.061	0.609	Green Entrepreneurial Orientation → Knowledge Management Orientation
0.000	14.155	0.049	0.692	Green Entrepreneurial Orientation → Market Orientation
0.000	7.816	0.072	0.564	Knowledge Management Orientation → GSCM
0.436	0.779	0.085	0.066	Market Orientation → GSCM

بر اساس نتایج این آزمون و مقادیر t به دست آمده که بالاتر از ۱/۹۶ می‌باشند، سه فرضیه میانجی تحقیق مورد تایید قرار گرفت.

جدول ۱۴: مقادیر t -value برای فرضیه‌های میانجی (از طریق آزمون سوبل)				
سطح معنی‌داری	مقادیر t	انحراف معیار	بار عاملی	
0.010	2.597	0.053	0.137	Green Entrepreneurial Orientation → GSCM → Economic performance
0.008	2.646	0.059	0.155	Green Entrepreneurial Orientation → GSCM → Environmental performance
0.000	6.554	0.052	0.343	Green Entrepreneurial Orientation → Knowledge Management Orientation → GSCM
0.441	0.771	0.059	0.046	Green Entrepreneurial Orientation → Market Orientation → GSCM

در شکل شماره (۲)، مدل ساختاری تحقیق به همراه مقادیر (T -Value) نشان داده شده است.



شکل (۲): مقدار ضریب (t) مدل ساختاری تحقیق

۵- نتیجه‌گیری

این تحقیق به منظور تحلیل تاثیر جهت‌گیری‌های کارآفرینانه سبز بر شیوه‌های مدیریت زنجیره تامین سبز در شرکت‌های شهرک صنعتی یزد به انجام رسید که نتایج یافته‌های تحلیلی آن به این شرح می‌باشد.

فرضیه اول: جهت‌گیری کارآفرینی سبز بر جهت‌گیری بازار تاثیر معنی‌داری دارد. با توجه به جدول مقدار قدرمطلق آماره t برابر ۱۴/۱۵۵ و بزرگ‌تر از مقدار ۱/۹۶ است، پس فرض صفر رد می‌شود. یعنی در سطح اطمینان ۹۵٪ جهت‌گیری کارآفرینی سبز بر جهت‌گیری بازار تاثیر معنی‌داری دارد. این نتایج همسو با یافته‌های تحقیق انجام گرفته توسط حبیب و همکاران (۲۰۲۱) می‌باشد.

فرضیه دوم: جهت‌گیری کارآفرینی سبز بر شیوه‌های مدیریت زنجیره تامین سبز تاثیر معنی‌داری دارد. با توجه به جدول مقدار قدرمطلق آماره t برابر ۲/۶۸۴ و بزرگ‌تر از مقدار ۱/۹۶ است، پس فرض صفر رد می‌شود یعنی در سطح اطمینان ۹۵٪ جهت‌گیری کارآفرینی سبز بر شیوه‌های مدیریت زنجیره تامین سبز تاثیر معنی‌داری دارد. این نتایج همسو با یافته‌های تحقیق انجام گرفته توسط برازون و همکاران (۲۰۲۱) و سیلوا، گومز، کراوالهو، و گرالسدس (۲۰۲۱) می‌باشد.

فرضیه سوم: جهت‌گیری کارآفرینی سبز بر جهت‌گیری مدیریت دانش تاثیر معنی‌داری دارد. با توجه به جدول

مقدار قدر مطلق آماره t برابر ۹/۹۲۵ و بزرگ‌تر از مقدار ۱/۹۶ است، پس فرض صفر رد می‌شود یعنی در سطح اطمینان ۹۵٪ جهت‌گیری کارآفرینی سبز بر جهت‌گیری مدیریت دانش تاثیر معناداری دارد. این نتایج همسو با تحقیق انجام گرفته توسط شهو و محمود^۲ (۲۰۱۴) بود

فرضیه چهارم: جهت‌گیری بازار بر شیوه‌های مدیریت زنجیره تامین سبز تاثیر معنی‌داری دارد. با توجه به جدول مقدار قدر مطلق آماره t برابر ۷/۷۹ و کوچک‌تر از مقدار ۱/۹۶ است، پس فرض صفر تایید می‌شود یعنی در سطح اطمینان ۹۵٪ جهت‌گیری بازار بر شیوه‌های مدیریت زنجیره تامین سبز تاثیر معنی‌داری ندارد. این نتایج مغایر با یافته‌های برخی از مطالعات بود. به عنوان نمونه عظیمی و همکاران (۱۴۰۰) نشان دادند، بازارگرایی بر راهکارهای مدیریت زنجیره تامین موثر بود، که همسو با یافته‌های مطالعه حاضر نمی‌باشد.

فرضیه پنجم: جهت‌گیری مدیریت دانش بر شیوه‌های مدیریت زنجیره تامین سبز تاثیر معناداری دارد. با توجه به جدول مقدار قدرمطلق آماره t برابر ۷/۸۱۶ و بزرگ‌تر از مقدار ۱/۹۶ است، پس فرض صفر رد می‌شود یعنی در سطح اطمینان ۹۵٪ جهت‌گیری مدیریت دانش بر شیوه‌های مدیریت زنجیره تامین سبز تاثیر معنی‌داری دارد. این نتایج همسو با یافته‌های تحقیق لیم، تسنگ، تان و بوی^۳ (۲۰۱۷) بود. به گونه‌ای که نتایج نشان داد مدیریت دانش

1- Silva, Gomes, Carvalho, & Gerales

2- Shehu & Mahmood

3- Lim, Tseng, Tan, & Bui

بر شیوه‌های مدیریت زنجیره تأمین سبز تاثیر معنی‌داری دارد.

فرضیه ششم: جهت‌گیری کارآفرینی سبز بر شیوه‌های مدیریت زنجیره تأمین سبز از طریق جهت‌گیری بازار تاثیر دارد. با توجه به جدول مقدار قدرمطلق آماره t برابر ۰/۷۷۱ و کوچک‌تر از مقدار ۱/۹۶ است، پس فرض صفر تایید می‌شود یعنی در سطح اطمینان ۹۵٪ جهت‌گیری کارآفرینی سبز بر شیوه‌های مدیریت زنجیره تأمین سبز از طریق جهت‌گیری بازار تاثیر ندارد. این نتایج مغایر با یافته‌های تحقیق انجام گرفته توسط حبیب و همکاران (۲۰۲۱) می‌باشد.

فرضیه هفتم: جهت‌گیری کارآفرینی سبز بر شیوه‌های مدیریت زنجیره تأمین سبز از طریق جهت‌گیری مدیریت دانش تاثیر دارد. با توجه به جدول مقدار قدرمطلق آماره t برابر ۶/۵۵۴ و بزرگ‌تر از مقدار ۱/۹۶ است، پس فرض صفر رد می‌شود؛ یعنی در سطح اطمینان ۹۵٪ جهت‌گیری کارآفرینی سبز بر شیوه‌های مدیریت زنجیره تأمین سبز از طریق جهت‌گیری مدیریت دانش تاثیر دارد. این نتایج

همسو با یافته‌های تحقیق انجام گرفته توسط حبیب و همکاران (۲۰۲۱) می‌باشد.

در پایان پیشنهاد می‌گردد: در خصوص تاثیر جهت‌گیری کارآفرینی سبز بر جهت‌گیری بازار، جهت‌گیری مدیریت دانش و شیوه‌های مدیریت زنجیره تأمین سبز پیشنهاد می‌شود تا در شرکت مورد مطالعه به شیوه‌های مدیریت سبز از قبیل: تحقیق و توسعه، نوآوری محصول و فرآیند، اولویت داده شده و در پاسخ به اقدامات رقبا از ابتکارهای سبز بهره گرفته شود. در خصوص تاثیر جهت‌گیری بازار و جهت‌گیری مدیریت دانش بر شیوه‌های مدیریت زنجیره تأمین سبز پیشنهاد می‌شود تا به طور منظم اطلاعات مشتریان را جمع‌آوری و مورد تحلیل قرار داده، بر اساس آن تصمیم‌گیری صورت گیرد. همچنین در حوزه مدیریت دانش شرکت پیشنهاد می‌شود تا قوانین روشنی برای طبقه‌بندی دانش محصولات تدوین و اجرا نموده و دانش موجود با سایر شرکت‌ها به اشتراک گذاشته شود.

منابع

- ۱- احمدی نژاد، سحر السادات؛ کریمی زارچی، محمد؛ و فتاحی، محمد رضا. (۱۳۹۹). انتخاب استراتژی تجاری مدیریت زنجیره تامین سبز با بکارگیری روش فرآیند تحلیل شبکه‌ای، *انسان و محیط زیست*. ۱۸ (۱)، ۲۱-۳۴.
- ۲- ایرجی راد، ارسلان؛ و پارسامهر، فاطمه. (۱۴۰۰). تاثیر جهت‌گیری استراتژیک حلقه بسته بر عملکرد با تبیین نقش میانجی مدیریت زنجیره تامین سبز در شرکت ایران خودرو. *فصلنامه مدیریت صنعتی*، ۵۵، ۳۴-۴۹.
- ۳- کیانی، مهرداد؛ مروتی شریف آبادی، علی؛ مفتاح زاده، الهام؛ و زمزم، فاطمه. (۱۴۰۱). طراحی مدل پویای عوامل موثر بر موفقیت مدیریت زنجیره تامین سبز، *بررسی‌های بازرگانی*، ۲۰ (۱۱۲)، ۲۵-۴۴.
- ۴- عظیمی، حسین؛ موسوی، سیده فاطمه؛ و خورشیدی، پوریا. (۱۴۰۰). بررسی تاثیر بازرگانی بر عملکرد محیط زیستی با نقش میانجی راهکارهای مدیریت زنجیره تامین سبز. *مجله پژوهش‌های محیط زیست*، ۲۳، ۲۳۲-۲۱۹.

- 5-Attia, A., & Essam Eldin, I. (2018). Organizational learning, knowledge management capability and supply chain management practices in the Saudi food industry. *Journal of Knowledge Management*, 22(6), 1217-1242.
- 6-Borazon, E. Q., Huang, Y. C., & Liu, J. M. (2022). Green market orientation and organizational performance in Taiwan's electric and electronic industry: the mediating role of green supply chain management capability. *Journal of Business & Industrial Marketing*, 37(7), 1475-1496.
- 7-Choi, D. (2014). Market orientation and green supply chain management implementation. *International Journal of Advanced Logistics*, 3(1-2), 1-9.
- 8-Copacino, W. C. (1997). *Supply chain management: The basics and beyond* (Vol. 1). CRC Press.
- 9-Dean, T. J., & McMullen, J. S. (2007). Toward a theory of sustainable entrepreneurship: Reducing environmental degradation through entrepreneurial action. *Journal of business venturing*, 22(1), 50-76.
- 10-Eltayeb, T. K., Zailani, S., & Ramayah, T. (2011). Green supply chain initiatives among certified

- companies in Malaysia and environmental sustainability: Investigating the outcomes. *Resources, conservation and recycling*, 55(5), 495-506.
- 11-Evers, N., Gliga, G., & Rialp-Criado, A. (2019). Strategic orientation pathways in international new ventures and born global firms—Towards a research agenda. *Journal of International Entrepreneurship*, 17, 287-304.
- 12-Faisal, A., Hermawan, A., & Arafah, W. (2018). The influence of strategic orientation on firm performance mediated by social media orientation at MSMEs. *International Journal of Science and Engineering Invention*, 22-to31.
- 13-Foo, P. Y., Lee, V. H., Tan, G. W. H., & Ooi, K. B. (2018). A gateway to realising sustainability performance via green supply chain management practices: A PLS-ANN approach. *Expert Systems with Applications*, 107, 1-14.
- 14-Frambach, R. T., Prabhu, J., & Verhallen, T. M. (2003). The influence of business strategy on new product activity: The role of market orientation. *International journal of research in marketing*, 20(4), 377-397.
- 15-Gast, J., Gundolf, K., & Cesinger, B. (2017). Doing business in a green way: A systematic review of the ecological sustainability entrepreneurship literature and future research directions. *Journal of cleaner production*, 147, 44-56.
- 16-Green, K. W., Zelbst, P. J., Meacham, J., & Bhadauria, V. S. (2012). Green supply chain management practices: impact on performance. *Supply chain management: an international journal*, 17(3), 290-305.
- 17-Groza, M. D., Zmich, L. J., & Rajabi, R. (2021). Organizational innovativeness and firm performance: Does sales management matter?. *Industrial marketing management*, 97, 10-20.
- 18-Habib, M. A., Bao, Y., & Ilmudeen, A. (2020). The impact of green entrepreneurial orientation, market orientation and green supply chain management practices on sustainable firm performance. *Cogent Business & Management*, 7(1), 1743616.
- 19-Hughes, P., Hodgkinson, I. R., Hughes, M., & Arshad, D. (2018). Explaining the entrepreneurial orientation–performance relationship in emerging economies: The intermediate roles of absorptive capacity and improvisation. *Asia Pacific Journal of Management*, 35, 1025-1053.
- 20-Islam, M. M., Perry, P., & Gill, S. (2021). Mapping environmentally sustainable practices in textiles, apparel and fashion industries: a systematic literature review. *Journal of Fashion Marketing and Management: An International Journal*, 25(2), 331-353.
- 21-Jiang, W., Chai, H., Shao, J., & Feng, T. (2018). Green entrepreneurial orientation for enhancing firm performance: A dynamic capability perspective. *Journal of cleaner production*, 198, 1311-1323.
- 22-Koberg, E., & Longoni, A. (2019). A systematic review of sustainable supply chain management in global supply chains. *Journal of cleaner production*, 207, 1084-1098.
- 23-Lim, M. K., Tseng, M. L., Tan, K. H., & Bui, T. D. (2017). Knowledge management in sustainable supply chain management: Improving performance through an interpretive structural modelling approach. *Journal of cleaner production*, 162, 806-816.
- 24-Lin, Y., Luo, J., Ieromonachou, P., Rong, K., & Huang, L. (2019). Strategic orientation of servitization in manufacturing firms and its impacts on firm performance. *Industrial Management & Data Systems*, 119(2), 292-316.
- 25-Luthra, S., Garg, D., & Haleem, A. (2015). Critical success factors of green supply chain management for achieving sustainability in Indian automobile industry. *Production Planning & Control*, 26(5), 339-362.
- 26-Majumdar, A., & Sinha, S. (2018). Modeling the barriers of green supply chain management in small and medium enterprises: A case of Indian clothing industry. *Management of Environmental Quality: An International Journal*, 29(6), 1110-1122.
- 27-Margahana, H., & Sugandini, D. (2022). Knowledge integration capability, innovativeness, and entrepreneurial orientation on business performance.
- 28-Niinimäki, K., Peters, G., Dahlbo, H., Perry, P., Rissanen, T., & Gwilt, A. (2020). The environmental

- price of fast fashion. *Nature Reviews Earth & Environment*, 1(4), 189-200.
- 29-Pacheco, D. F., Dean, T. J., & Payne, D. S. (2010). Escaping the green prison: Entrepreneurship and the creation of opportunities for sustainable development. *Journal of Business Venturing*, 25(5), 464-480.
- 30-Paulraj, A. (2011). Understanding the relationships between internal resources and capabilities, sustainable supply management and organizational sustainability. *Journal of Supply Chain Management*, 47(1), 19-37.
- 31-Sahi, G. K., Gupta, M. C., & Cheng, T. C. E. (2020). The effects of strategic orientation on operational ambidexterity: A study of indian SMEs in the industry 4.0 era. *International Journal of Production Economics*, 220, 107395.
- 32-Shehu, A. M., & Mahmood, R. (2014). Market orientation, Knowledge management and Entrepreneurial orientation as predictors of SME performance: Data screening and Preliminary Analysis. In *Information and Knowledge Management* (Vol. 4, No. 7, pp. 12-23).
- 33-Silva, G. M., Gomes, P. J., Carvalho, H., & Geraldles, V. (2021). Sustainable development in small and medium enterprises: The role of entrepreneurial orientation in supply chain management. *Business Strategy and the Environment*, 30(8), 3804-3820.
- 34-Tseng, S. M. (2014). The impact of knowledge management capabilities and supplier relationship management on corporate performance. *International Journal of Production Economics*, 154, 39-47.
- 35-Tseng, M. L., Islam, M. S., Karia, N., Fauzi, F. A., & Afrin, S. (2019). A literature review on green supply chain management: Trends and future challenges. *Resources, Conservation and Recycling*, 141, 145-162.
- 36-Zaim, H., Keceli, Y., Jaradat, A., & Kastrati, S. (2018). The effects of knowledge management processes on human resource management: Mediating role of knowledge utilization. *Journal of Science and Technology Policy Management*, 9(3), 310-328.
- 37-Zhu, Q., Sarkis, J., & Lai, K. H. (2007). Green supply chain management: pressures, practices and performance within the Chinese automobile industry. *Journal of cleaner production*, 15(11-12), 1041-1052.

©Authors, Published by Journal of Intelligent Knowledge Exploration and Processing. This is an open-access paper distributed under the CC BY (license <http://creativecommons.org/licenses/by/4.0/>).



components and their importance. This study is the first attempt to identify the effects of GSCM in the form of strategic orientation.

2- Theoretical Foundations

The orientation of green entrepreneurship is defined as a perspective to change society through pollution prevention and reducing environmental effects. The orientation of green entrepreneurship is undoubtedly a combination of some characteristics of entrepreneurship as a decision-making situation in the direction of the decision-making process. Market orientation is defined as strategies and operations of the company that are used to meet the current and future use of the product and the market. Market orientation can also be defined as the ability to identify and meet customers' needs. Knowledge management orientation helps the organization to maintain proper communication with up-to-date information and also to develop and manage good relationships with its suppliers to grow business. Supply Chain Management is the coordination of production, inventory, location and transportation between participants in a supply chain, which is used to achieve the best combination of responsiveness and efficiency for success in the market. GSCM from the perspective of the product life cycle includes all stages of raw materials, product design and manufacturing, product sales and transportation, product use and product recycling. By using Supply Chain Management and green technology, the company can reduce negative environmental impacts and achieve optimal use of resources and energy. Strategic orientation is one of the guiding principles affecting marketing and choosing strategies for a company's activities, which reflects the strategic trends implemented by the company to build appropriate behaviors and leads to better performance. It is formed based on the company's beliefs of doing business through a wide set of values and fundamental beliefs. The background of the research is stated in table number (1).

3- Research Methodology

The present study is analytical and cross-sectional in terms of time, which was carried out

as a survey on the employees of Yazd Industrial City. In this research, the statistical population is 300 employees of Yazd Industrial City companies, using the Cochran's formula, the final sample included 168 participants. After distributing the questionnaires, 150 questionnaires were collected. The simple random sampling method was adopted. In the data analysis of this research, descriptive statistics indicators such as frequency, percentage, average and standard deviation were used to describe the variables. Analytical tests were used to examine the hypotheses. Also the appropriate method of structural equation modeling and PLS and SPSS software have been used.

4- Conclusion

Based on the reviews, there are limited studies that have systematically examined strategic ways to achieve superior performance of the company while ensuring sustainability, especially in terms of business and production. The sustainable growth of industries in the face of significant negative impact on environmental sustainability is a big challenge, and manufacturers are trying to implement GSCM and improve company performance. Therefore, there are few studies on the relationship between strategic orientation and performance of the company. Accordingly, this article investigated how market orientations and knowledge management play a mediating role in the impact of green entrepreneurship orientation on GSCM. The results of which are as follows: green entrepreneurship orientation has a significant impact on market orientation, GSCM practices, and knowledge management orientation. Market orientation has a significant impact on GSCM practices. Knowledge management orientation has a significant impact on GSCM. Green entrepreneurship orientation has an impact on GSCM through market orientation. The orientation of green entrepreneurship has an effect on GSCM practices through the knowledge management orientation.

Keywords: Strategic Orientations, Supply Chain Management, Green Supply Chain (GSC), Sustainable Performance, Green Entrepreneurship.

Research Article

Analyzing the Impact of Green Entrepreneurial Orientations on Green Supply Chain Management Practices

Doi: 10.30508/KDIP.2023.379963.1058

Masoud Bakhtiari Azadeh (Corresponding Author)¹ | Bahareh Shahriari²

Abstract

1- Introduction

Yazd province is one of the most industrial regions of Iran and many days of the year, the weather is highly polluted. However, it is not easy to implement Green Supply Chain Management (GSCM) without motivation, pressure, top management support, green infrastructure, green supply and green information system facilities. Government's force, organizational culture, green knowledge, technological and financial capability, socioeconomic benefits from green practices might be the background for adopting GSCM in the wire and cable industry. Although some companies have accepted ISO 9000 and ISO 14001 certificates, many industries are lagging behind in the implementation of GSCM. Hence, organizations should take strategic measures in adopting practices to improve sustainable performance. Strategic orientation is one of effective principles in marketing and choosing strategies adopted by a company which reflects the strategic trends implemented by the company to build appropriate behaviors and leads to better performance, and Strategic orientation is based on the company's beliefs about doing business and formed through a wide set of values and fundamental beliefs. Therefore, the general view of strategic orientation in the implementation of the best GSCM practices regarding sustainability performance lacks theoretical and practical literature. In addition, previous researches in the field of GSCM have mainly focused on simple direct effect analysis, or identification of

1- Department of Accounting, Payame Noor University (Tehran Center), Tehran, Iran

2- Department of Management, Payame Noor University (Tehran Center), Tehran, Iran

based on fuzzy logic reasoning strategy for edge detection in digital images without determining a threshold value. The proposed method starts by segmenting the image into different areas, using a 3x3 Binary matrix, so that the edge pixels are written in the range or different values from each other. This method causes a stable effect on the smoothness of the lines for straight lines and rotation for curved lines. In another study, Flores Vidal and et al. edge detection operation has been done with fuzzy clustering method; traditionally, it is observed that the edge detection process requires a final step known as scaling. This is done to make decision pixel-by-pixel as to whether or not it should be selected as the final edge. Although many edge detection methods have been proposed for grayscale, color, and multispectral images, they still face problems when extracting edge features from hyper spectral images that contain a large number of bands with very narrow gaps in the spectral domain.

3- Research Methodology

The algorithm starts by dividing the pixels of the HI image into 3x3 areas (masks), and then the image is mapped to a range of values by the fuzzy inference system and edge detection is performed. The main steps in edge detection by fuzzy method are: fuzzification of images, membership values modification and de-fuzzification if necessary. Due to the 8-bit input image and the output image obtained after de-phasing, the gray levels are always between 0 and 255 values. Fuzzy sets are produced to show the intensity of each variable. These sets depend on the linguistic variables of black color, white color and edge. Also the membership function considered for input and output fuzzy sets is triangular. In the inference stage, Min and Max functions are used for And and Or operations. The de-fuzzification and

inference procedure is Mamdani method. Fuzzy and non-fuzzy steps are due to the failure of fuzzy hardware processing. Therefore, encoding the image data and decoding the results are steps that enable image processing with fuzzy techniques. The main power in the middle step is to modify the fuzzy membership values. Then the image data is converted from the gray level plane to the fuzzy membership plane. Appropriate fuzzy techniques modify the fuzzy membership values and cluster the image.

4- Conclusion

Edge detection is one of the important techniques used for image segmentation. Image segmentation is still an important issue even after four decades. In addition to the traditional methods in edge detection, in recent years, various methods and techniques have been proposed for edge detection. Most of the presented methods are based on the combination of traditional methods in edge detection. The method proposed in this research shows better results than some common methods due to maintaining edge continuity and removing noise. In this way, using the fuzzy concept, a new method for edge detection has been presented, in which the image is considered as a fuzzy set and the pixels are considered elements of the fuzzy set. The proposed fuzzy method converts the color image into a clustered image, finally an edge detector operates on the clustered image to obtain the edge image. As the proposed operator does not perform the darkening operation on the image, it also detects the double edges less. Generally, real images include strong and weak edges. The proposed method gives both strong and weak edges with different edge strength using upper and lower threshold limits.

Keywords: Image Processing, Edge Detection, Fuzzy Theory, Fuzzy Clustering.

Research Article

Edge Detection of Digital Images Using Optimized Fuzzy Clustering

Doi: 10.30508/KDIP.2023.368960.1054

Moheb Ali Soufi (Corresponding Author)¹ | Alireza Nasiri²

Abstract

1- Introduction

Image is one of the biggest information sources of human intelligence activities. With the passage of time, the need for image processing has also increased day by day. Edge detection is one of the operations used in image analysis as well as necessary pre-processing in image segmentation. An edge is the boundary between two areas of the image. Edge detection operations are essential in different areas of image processing such as; Feature extraction, feature recognition, pattern recognition, etc. All edge detection methods are classified under two groups: Gradient and Laplace. Among the edge detection methods, the Laplacian, Roberts, Sobel, Prewitt, Canny and fuzzy logic methods are the most widely used ones. The presence of noise and blur image are among the problems you encounter in these techniques. Clustering is one of the branches of unsupervised learning and is an automatic process during which samples are divided into categories whose members are similar to each other.

2- Theoretical Foundations

One of the edge detection methods is the use of classic operators such as linear filters. Although these methods have less power in edge detection, they are used more due to the ease of implementation. Bahagamabati and et al. have proposed a new method

1- MA graduate in Artificial Intelligence and Robotics, Hormozgan University, Hormozgan, Iran

2- Assistant Professor of Electrical Department, Hormozgan University, Hormozgan, Iran

that precede these activities and determine them. Consumer innovation: the current era is the era of transformation and unprecedented and tight competition between international economic companies. Conditions are constantly changing due to the disappearance of borders in marketing, market fragmentation, shortening of product life, rapid changes and methods of using customer services and increasing consumer awareness. Brand performance: Brand performance is a criterion for measuring brand success in the market. Online brand: An online brand is defined as a brand management technique which uses the Internet as a mediator to place the brand name in the market. Internet marketing: It is a transaction in which the purchase or sale of goods or services takes place through the Internet and leads to the import or export of goods or services. Brand performance in the online space: The success of a business undoubtedly depends on the brand performance of that business. The need to measure the organization's performance from different aspects and according to different levels has often been considered in the marketing literature as a dependent variable, therefore there is a view of evaluating the performance through the products and services provided by the organization. In other words, often in the brand issues, two main questions arise in the mind: "What factors make the brand strong?" and "How to build a strong brand?" To answer these questions, the broad concept of Brand Performance has been introduced. Therefore, the managers are better equipped by being aware of the dimensions and characteristics of brand performance and can apply more effective strategies for the brand. Interaction with the online brand: the seller tries to attract customers mainly by using online services and virtual networks based on customer interaction in the virtual and internet space and tries to introduce his products and services with the lowest price and the highest facilities.

3- Research Methodology

This research sets to investigate the relationship between factors affecting consumer innovation

and the performance of online brands with regard to the mediating role of customer interaction with the online brand (the case study of Ocala store in Mashhad). The statistical population includes all the citizens of Mashhad who buy online from Ocala, according to the characteristics of the statistical population and the available simple random sampling method and the use of Morgan method, 384 cases were considered and examined.

4- Conclusion

In order to create effective and long-term relationships, it is necessary to know the effective bottom layers in consumer relations, and insufficient understanding of emotional links destroys all the company's marketing costs. On the other hand, creating emotional bonds turns customers into fierce guardians of the brand, who become worried and anxious when they are separated from it, and when they have the brand, they support it like fierce defenders. Today, many companies in the field of online business have come to the conclusion that the most valuable asset of a company in order to improve the marketing process is the brand and its performance. Nevertheless, the acceptance of brands in the online environment is one of the most important topics in the field of information systems and brands, which has received less attention. Despite the many relative advantages of online business, it has not been well received in Iran and some developing countries. One of the reasons for not buying online is not accepting risk. Trust in the brand, which is considered as one of the cognitive and effective components in the relationship between the customer and the brand, means the customer's belief in the reliability, honesty of the brand, and its loyalty to the promises. In this way, despite the competition between Iranian companies in attracting customers, such companies still do not use the Internet, especially social networks, to attract customers.

Keywords: Consumer Innovation, Brand Performance, Online Brands, Customer Interaction, Online Brand Interaction.

Research Article

Consumer Innovation and Performance of Online Brands Considering the Mediating Role of Customer Interaction with Online Brand

(Case Study of Ocala Store in Mashhad)

Doi: 10.30508/KDIP.2022.366246.1052Ali Mirzaei Rad¹ | Mohammad Hossein Homayouni Rad² | Saeedeh Babajani Mohammadi³

Abstract

1- Introduction

Today, what is important for customers is the use of products and services that save time and increase customer comfort. The desire to receive services from organizations that care about customer comfort is due to the customer's desire for a more comfortable lifestyle, and the ease of using services is a key indicator for them. The customer's understanding of convenient use of the service, affects the customer's satisfaction and evaluation of the service, perceived service credibility, his trust and loyalty to the brand. The importance of each feature of ease of understanding and use of services is very variable in terms of customer quality; customers who give value to saving time, he will give more credit to use of services' ease.

Currently, customers expect the brands or goods they buy to satisfy them; But satisfaction is not a sufficient condition to create a continuous relationship with the brand. Customers are looking for connections beyond satisfaction; Bonds that are based on emotional dependency. On the other hand, Internet users are gradually forming business relationships that have traditionally been formed by marketers. In most developed countries of the world, social networks are widely used and have covered almost all aspects of people's lives. The organizers of these networks have been able to make the best use of these tools and the producers of products and services have been able to strengthen trust and loyalty to the brand in customers based on the information they receive from these networks.

2- Theoretical Foundations

Consumer behavior: the activities of people who are directly involved in the provision and use of economic goods and services, including the decision-making processes

1- Master of Business Administration, Ferdows Institute of Higher Education, Mashhad, Iran.

2- Assistant Professor of Department of Management, Sanabad Institute of Higher Education, Golbahar, Iran.

3- Assistant Professor of Department of Management, Ferdows Institute of Higher Education, Mashhad, Iran.

Beside all these metrics that are based on the centrality of the special vector, there are a series of centrality metrics that are based on the influence of penetration based on the independent cascade model; According to this model, each member is likely to influence the target member. To propagate this effect, the member must be active and only has one chance to affect another member. When a member tries to influence another, this member is considered active and this procedure will be repeated. The whole process ends when there is neither an active member nor a chance to influence.

A social network can be displayed in the form of a graph in which nodes, network members and edges display personal relationships between members. Sometimes edges have a weight that expresses the strength of each relationship. A less common feature is node labels, which are used to display different network features. In this work, the labels indicate the resistance of the node to influence, while the weight of the edges indicates the strength of influence applied from one member to another. A weighted and labeled social network can be defined more formally in the form of an influence graph.

3- Research Methodology

An example of a graph with a high linear concentration threshold. Graph 'b' shows two large cores instead of one. However, it is highly concentrated in the center, there are nodes (members) in the surface that can be reached by only one or two cores, and accordingly our measure returns a lower value than the first graph. Finally, when the core of both graphs 'c' and 'd' are at a lower level, the initial activation energy will probably not be enough to affect

the peripheral members, which is the result of measuring concentration at lower energy levels.

4- Conclusion

This article has concentrated on the comparison of centrality measures that were derived from the LTM. It was investigated in different social networks. Linear threshold ranking is introduced as: A new centrality measure based on the linear threshold model for each member has been presented. It has the capability of presenting influence diffusion. This metric is compared with three other centrality metrics (Katz centrality, page rank, and independent cascade ranking) based on penetration and influence in four real networks, two large networks (one directed, the other undirected) and two small networks (one directed, the other undirected). In general, the findings show that the true value of correlation results in Spearman's coefficient is higher than Kendall's coefficient. In addition, linear threshold ranking with independent cascade ranking yields the highest deviance. In addition to the centrality measure, a centralization measure called linear threshold centralization has been introduced in this article; which corresponded to the number of members that could be influenced by nodes belonging to the core of the network. It is proven that this metric will be very useful for networks with a large number of members. Also, this metric was compared with the average clustering coefficient and it was concluded that each of these metrics creates different centralization and is used to extract different information from the network.

Keywords: Centrality, Linear Threshold Model, Independent Cascade Model, Spread of Influence, Social Network.

Research Article

Centrality in Social Networks Based on Linear Threshold Model

Doi: 10.30508/KDIP.2022.361686.1049

Ali Sarabadani¹ | Kheirollah Rahseparfard (corresponding author)²

Abstract

1- Introduction

A social network can be shown as a graph where the nodes represent the members of the network and the edges represent the relationships between the members, sometimes the edges are weighted to indicate the strength of each member's relationship. Meanwhile, the centrality measures show the structural connection of each member with this social network. In recent years, due to the increase of Internet users and the creation of various and complex social networks, the emergence of more complex centrality measures can be clearly observed. Accordingly, as per the large volume of these networks, i.e., the large number of members and the relationships between them, it is necessary to calculate the metrics in a practical way. One of the most well-known general models for diffusion of influence is the Linear Threshold Model (LTM), which was based on a series of ideas of collective behavior, and the independent cascade model has been proposed for the subject of marketing.

2- Theoretical Foundations

In this section, the known centrality measures have been discussed. One of the most widely used measures among related indicators is eigenvector centrality, which indicates that a member in a group is important if it has important neighbors linked to it or if it has many members as neighbors.

1- PhD Candidate in Information Technology Engineering / E-Commerce, University of Qom, Qom, Iran

2- Assistant Professor, Department of Computer Engineering and Information Technology, University of Qom, Qom, Iran

design new sophisticated fraud attacks on a daily basis, thus requiring researchers to continuously develop new fraud detection techniques.

2- Theoretical Foundations

The specialized literature and solutions on the market in the last decade have focused on collecting significant amounts of data about transactions (and user behavior) and modifying the algorithms used to detect fraud. In this regard, EU legislation (e.g., PSD2) has been adopted to impose obligations on beneficiaries to detect theft. However, on the one hand, the law provides a high-level description of this legal requirement, and on the other hand, available solutions are diverse in terms of data collected and the effort to collect data for production.

Fraud in various forms increases every year, which leads to significant financial losses. According to Microsoft Computing Safety Index (MCSI), in 2014 the annual impact of phishing and various forms of identity fraud around the world can be estimated at up to US\$5 billion, while the cost of repairing damage to an individual's online reputation can exceed US\$6 billion, or an estimated average of US\$ 632 per loss. Fraud detection involves activities that prevents money or property from being obtained through fraudulent misrepresentation. In the banking sector, fraud involves check forgery and using stolen credit cards. It may involve exaggerated losses or creating unfortunate events such as accidents for the sole purpose of payment.

3- Research Methodology

Logistic regression refers to a supervised learning method used with deterministic decisions. This means that if the transaction is carried out, the obtained results can be considered as fraud or non-fraud. This approach uses a cause and effect relationship to create an organized data set. Regression analysis technique when

used in fraud detection is more complex due to multiple variables and data set size. This model (algorithm) predicts whether or not new transactions will be marked as fraudulent. These models are more accurate than the models that merchants themselves have adopted, but generic models are still applicable. Decision tree algorithms are used to classify unusual activities in a transaction from an authorized user. These algorithms have limitations in the dataset used in fraud classification. Decision tree algorithms are used in classification or regression extrapolation modeling problems. They are basically a set of rules that can be utilized to apply for fraud of customers.

4- Conclusion

Fraudulent activity has increased during the Corona Virus period, and organizations with traditional fraud detection methods have proven to have impressive results; the human eye does not always detect anomalies in the banking system. Financial fraud is one of the main problems that undermines the customer's trust and also causes economic losses to banks and financial institutions. In recent years, with the spread of fraud, financial institutions have been looking for a right solution to deal with fraud. Due to the advanced and various changes in fraud methods, extensive research has been done to detect them. Banks seek to reduce huge losses through fraud detection and prevention systems. Many advanced fraud technologies are used to detect and prevent fraudulent Internet banking transactions. However, they have no effective identification mechanism to identify legitimate users and track their illegal activities.

Keywords: Bank Fraud, Electronic Banking, Fraud, Fraudulent Transactions, Fake Transactions, Bank Fraud Detection.

Research Article

An Overview of Bank Fraud Detection Methods Using Artificial Intelligence

Doi: 10.30508/kdip.2022.349333.1044

Seyyed Majid Akbar Alsadat (Corresponding Author)¹ | Babak Ismailpour²

Abstract

1- Introduction

Fraud, due to its destructive effects is one of the most threatening problems that every human society is dealing with. It refers to the intentional use of false information to defraud the money or properties of another person or organization. In most cases, for the fraudster, the unmet needs are endless and different for various people. When a person has unmet needs and limited resources, he recognizes the opportunity to commit fraud. Therefore fraud detection methods must constantly evolve as fraudsters become more capable at developing techniques that bypass bank security systems and learn how to convince unsuspecting people to hand over their money. Fraud detection for online banking is a very important field of study, as cybercriminals

1- Department of Artificial Intelligence and Robotics, Islamic Azad University (Quchan Branch), Quchan, Iran

2- Assistant Professor of Computer Department, Islamic Azad University (Mashhad Branch), Mashhad, Iran

process in cloud computing, various methods and algorithms have been presented in order to achieve acceptable scheduling in cloud computing. In all these algorithms, an optimal solution has been examined to obtain a suitable schedule. Each of these implemented methods has their advantages and disadvantages, and a general solution for improving all parameters in the cloud space has not been proposed. In this paper, the dynamic management of resources has been done based on the effective combination of virtual machines and it has been implemented by the proposed algorithm according to the Service Level Agreement (SLA). The difference between this method and the other methods is that it shows the improvement of time parameters, speed and accuracy and convergence. Experimental results show that the presented method performs better in terms of energy consumption, number of Virtual Machine Migrations, and Service Level Agreement Violations.

3- Research Methodology

The main purpose of task scheduling is to prioritize and allocate tasks to free resources in such a way that maximum tasks are being performed in parallel. Other scheduling goals include: load balancing, service quality, economic principle, the best execution (running) time and system throughput. Different scheduling methods and techniques are used to assign tasks to different nodes and virtual machines in cloud environments. The presented scheduling algorithms can be

classified based on the improvement parameters in the cloud environment. Scheduling algorithms are based on load balancing, cost, priority, overall completion time reduction, energy consumption, compatibility, etc. In this section, there are methods in the field of task scheduling and examining the advantages and disadvantages of these methods.

4- Conclusion

Task scheduling plays a key role in cloud computing systems, and this type of scheduling can not be done only based on a single criterion, but a large number of rules and conditions must be considered as an agreement between users and cloud providers. In fact, this agreement is the quality of service that users expect from providers. Cloud data centers must not only perform these massive tasks but also satisfy the multiple needs of different users. Scheduling in distributed computing systems is an important issue and various scheduling algorithms have been presented in this field. Parallelization method can be used to reduce the computational time complexity of the presented genetic algorithm. Parallel processing is one of the most important advantages of genetic algorithm. This means that in this method, instead of a variable, we grow a population towards the optimal point at a time. Therefore, the rate of convergence of the method will be very high.

Keywords: Virtual Machine Locating, Cloud Data Centers, Task Scheduling.

Research Article

Task Scheduling in Resource Balance for Virtual Machine Locating in Cloud Data Centers Using an Improved Genetic Algorithm

Doi: 10.30508/kdip.2023.346375.1041

Mohsen Azadeh¹ | Ahmad Baratian² | Hamid Tabatabaee (Corresponding Author)³

Abstract

1- Introduction

The performance of computer systems depends on several factors. One of these factors is the balancing of requests. This factor depends on the amount of work assigned to the system in a certain period of time. Hence, computers should be managed based on the principles of task priority. Task scheduling is a method in cloud systems that is responsible for distributing computing tasks among virtual resources. Cloud data centers provide information based on the needs of users, regardless of their geographical location. Cloud data centers have been focused on timely delivery of services, reliability, and providing security in sending and receiving data, fault tolerance in service provision, stability and creating scalable infrastructure to host web-based application services. In general, cloud computing is a model that allows users' access based on the type of demands that have informative and computing resources. With the increasing rate of using this technology around the world, as well as the expansion of cloud computing, virtualization has been used as the best solution in this field. By using virtualization technology, instead of buying several servers for data storage, we can save our data on cloud data centers and pay only for the amount of our use, and if needed, we can use more resources.

2- Theoretical Foundations

Virtual machines locating is of particular importance due to the increasing complexity of cloud systems. Considering the high importance of the scheduling

1-2-3- Computer Engineering Department, Islamic Azad University (Mashhad Branch), Mashhad, Iran

Editor-in-Chief's Message

In the Name of God

Regarding the necessity of applying knowledge management in the organizational structures and discourse of knowledge management in governmental organizations and even in private sectors, and considering that the **"Intelligent Knowledge Exploration and Processing"** Journal has published the sixth issue, we witness potential positive aspects with regard to paying attention to knowledge assets and transfer of experiences and lessons obtained in organizational processes and inter-organizational communication, which is also subject to the judgment of professors and researchers in the upcoming issue. Transforming these potential strengths into actual advantages helps flourishing scientific, organizational and administrative environments, improving efficiency and developing innovation. Therefore, the realization of these advantages requires careful planning, the use of collective wisdom, the capabilities of scientific centers, and the participation of the employees and beneficiaries of the institutions. The beating heart of this process is the transformation of capabilities to their realizations, paying attention and benefiting from the competencies of management and knowledge management specialists who have been trained for years by the country's first-rank universities and higher education centers as they are ready to be employed. Accordingly, the staff of the journal are open to receiving scientific articles and valuable opinions of educated academic researchers, we wish we can take a small and effective step in line with our scientific mission with sincere effort and seriousness. In the end, it necessary that we express our gratitude to the continuous efforts of the managers, faculty members and academics who helped **Ferdows Institute of Higher Education**, also referees who evaluated and reviewed the received articles of the journal, as well as contributors to the articles.



Contents

Editor-in-Chief's Note	3
Task Scheduling in Resource Balance for Virtual Machine Locating in Cloud Data Centers Using an Improved Genetic Algorithm	4
An Overview of Bank Fraud Detection Methods Using Artificial Intelligence	6
Centrality in Social Networks Based on Linear Threshold Model	8
Consumer Innovation and Performance of Online Brands Considering the Mediating Role of Customer Interaction with Online Brand	10
Edge Detection of Digital Images Using Optimized Fuzzy Clustering	12
Analyzing the Impact of Green Entrepreneurial Orientations on Green Supply Chain Management Practices	14

VOL 2- ISSUE 6- Fall 2022

Print ISSN: 27833607-

Online ISSN: 27833615-

■ **■ Concessionaire: Ferdows Institute of Higher Education**

Director-in-Charge: Hamid Tabatabaee, Assistant Professor

Editor-in-Chief: Ebrahim Mahmoudzadeh, Assistant Professor

Deputy Editor: Saeedeh Babajani Mohammadi, Assistant Professor

Internal Manager: Sakineh Ghasemi, Engineer

■ **■ Editorial Board**

Mahmoud Moghavvemi

Professor, Department of Electrical Engineering, Faculty of Engineering, Universiti Malaya, Malaysia.

Mohamed Othman

Professor, Department of Communication Technology and Network, Faculty of Computer Science and Information Technology, Universiti Putra Malaysia (UPM).

Raja Syamsul Azmir b. Raja Abdullah

Professor, Department of Computer and Communication Systems Engineering, Faculty of Engineering, Universiti Putra Malaysia (UPM).

Logeswaran Rajasvaran

Professor, School of Computing, Asia-Pacific University of Technology and Innovation, Malaysia.

Bahman Moghimi

Professor, Faculty of Management and Economics, University of Georgia, Tbilisi, Georgia.

Mehrdad Jalali

Associate Professor and Scientist, Karlsruhe Institute of Technology (KIT), Germany.

Peyman Khavan

Professor, Ghom University of Technology - President of the Iranian Knowledge Management Scientific Association, Iran.

Reza Hasnavi Atashgah

Professor, Faculty of Industrial Engineering, Malek Ashtar University, Tehran, Iran.

Amir Masoud Rahmani

Professor, Faculty of Mechanics, Electrical and Computer Science, Islamic Azad University, Tehran, Iran.

Mahmoud Rezaei Roknabadi

Professor, member of the board of trustees of Ferdows Institute of Higher Education and member of the academic staff of Ferdowsi University of Mashhad, Iran.

Ebrahim Mahmoudzadeh

Professor, Malek Ashtar University of Technology, Tehran, Iran.

Ali Moeini

Professor, Faculty of Engineering, University of Tehran, Iran.

Mohammad Mehr-Aein

Professor, Faculty of Administrative and Economic Sciences, Ferdowsi University of Mashhad, Iran.

Amin Jajarmi

Associate Professor, Department of Electrical Engineering, University of Bojnord, Iran.

Javad Hamidzadeh

Associate Professor, Faculty of Computer and Information Technology, Sajjad University of Technology, Mashhad, Iran.

Abbas Ali Rezaei

Associate Professor, Payame Noor University of Mashhad, Iran.

Morteza Faraji

Associate Professor, Member of the Board of Trustees of Ferdows Institute of Higher Education and member of the Academic Tehran National Defense University, Tehran, Iran.

Mohammad Hossein Moattar

Associate Professor, Islamic Azad University of Mashhad, Iran.

Saeedeh Babajani Mohammadi

Assistant Professor, Department of Management, Ferdows Institute of Higher Education, Mashhad, Iran.

Alireza Rouhani Manesh

Assistant Professor, Department of Electrical Engineering, Faculty of Engineering, University of Neyshabur, Iran.

Mohammad Hadi Zahedi

Assistant Professor, Khajeh Nasir Toosi University of Technology, Tehran, Iran.

Seyed Kazem Shokfeta

Assistant Professor, Department of Computer Engineering, Shandiz Institute of Higher Education, Mashhad, Iran.

Hamid Tabatabaee

Assistant Professor, Department of Computer Engineering, Islamic Azad University of Mashhad, Iran.

Mojtaba Kafashan Kakhki

Assistant Professor, Department of Information Science and Knowledge, Ferdowsi University of Mashhad, Iran.

Abbas Mehdizadeh

Assistant Professor, Department of Computer, Ferdows Institute of Higher Education, Mashhad, Iran.

Persian Editor: Saeedeh Babajani Mohammadi

English Editor: Abbas Mehdizadeh

Headline and Cover Design: Mohammad Mohsen Khezri

Page Layout and Grid Design: Nima Malekzadeh

Magazine Expert: Ahad Fani Maleki

Address: Ferdows Institute of Higher Education, Kolahdouz 30, Shahid Kolahdouz Blvd., Mashhad, Iran.

Website: www.kdip.ir

Phone: +98 051337138011- ext. 703 and 716,051-5-372911114

Email: journal.kdip@gmail.com