

PSYCHOLOGY

Paper 9990/12
Paper 1 Approaches, Issues and
Debates

Key messages

Candidates need to know all components of every core study as listed in the syllabus. Questions can be asked about any part of a core study.

Candidates need to read the whole question carefully to ensure that their responses are fulfilling the demands of each one. For example, the question may require data, a named issue to be included or to relate back to a previous answer. To achieve full marks, these need to be correctly present in their responses. The essay (final question) requires four evaluation points to be in depth (two strengths and two weaknesses) with at least one of these about the named issue. In depth tends to mean having two examples of a particular concept or to support an evaluative point. Credit is limited if the named issue is omitted or just described.

Candidates need to be careful about how they are presenting the results of studies. For example, they need to know if the results are about how many participants performed a task correctly or on how many trials the participant was correct. This can have a large impact on the interpretation of results and whether a response can gain credit. In addition, there must never be any value judgements with results. Results are the factual presentation of data without any interpretation. Therefore, using words like 'better' or 'worse' are not appropriate in relation to the presentation of results.

Candidates also need to engage with any stimulus material presented in a question (for example, a novel situation) to ensure they can access all available marks. In addition, when a question refers to 'in this study' the answer requires contextualisation with an explicit example from that study.

Candidates also need to know the set procedure of studies in the order presented in the original journal article. Questions can be based around just *part* of a procedure and the candidate must be able to produce an answer that is directed and concise rather than writing about the whole of the procedure, otherwise a candidate may run out of time for other questions.

Candidates should be able to give full definitions of terms listed in the syllabus and provide full assumptions for all four approaches.

There is enough time for answers to be planned to ensure that the response given by a candidate is focused on the demands of each question. This is a crucial skill to develop as some candidates appear to have good knowledge of a study but do not apply this effectively to the question(s) set.

General comments

The marks achieved by the candidates sitting this examination covered a wide spread of possible marks. However, three-quarters of the candidates scored 25 marks or above (and one-quarter scored 41 marks or above). Some candidates provided a range of excellent answers to many of the questions and could explain psychological terminology well providing evidence that they were prepared for the examination.

Stronger overall responses followed the demands of each question with explicit use of psychological terminology and logical, well-planned answers in evidence. Appropriate examples were used from studies when the question expected it and there was evidence of candidates being able to apply their knowledge to real-world behaviours in terms of what and how.

There were several blank responses in this series (twelve of the items had blank responses). As positive marking is used, candidates should attempt all questions even if they are unsure of the response they are providing.

Performance on **Questions 3, 6, 9 and 10** had the strongest correlation with a candidate's overall score.

Comments on specific questions

Question 1

- (a) The majority of responses could correctly identify one of the other approaching figures. Common errors included 'parents' or repeating 'friend' from the question. It is essential for candidates to read the question carefully to ensure they are focused on what it is asking, especially when information is provided that cannot be used in a response.
- (b) Stronger responses could clearly outline what is meant by the term empathy. These tended to be able to cover both components of the definition. There were some responses that were tautological (e.g., empathy is about being empathic) which could not be given any credit. Also, some candidates mixed up empathy and sympathy.
- (c) Stronger responses could clearly provide a full conclusion to the study by Perry *et al.* The most popular choice was to focus on the role of oxytocin in high- and low-empathy people and how that affected personal space choices. A minority of responses presented a result rather than a conclusion. A result is based on factual data collected in the study whereas a conclusion is the general 'meaning' of the data. It is important for candidates to be able to differentiate between a result and a conclusion.

Question 2

- (a) A significant minority of responses could provide a full outline of what happened during 'body scan'. Popular choices included being guided through the body and seeing the body as a whole. However, the majority of responses focused on the brain scan rather than the body scan. This is a good example of candidates needing to read the question carefully. In the study by Hölzel *et al.* it was only the brain that got scanned, not the entire body. This question had the highest rate of blank responses on the paper.
- (b) A majority of responses could provide an appropriate weakness and then contextualise it via an example from the study. Common choices included ethical issues, sample generalisability and extraneous factors potentially affecting grey matter. However, there were some responses focused on the study being a 'laboratory experiment'. It is not. The brain scans took place in a controlled laboratory setting as they have to, but the study was not exclusively conducted in a laboratory. It is important for candidates to know which methods are used for all 12 core studies. In addition, some candidates questioned the replicability of the study as participants might have done different activities at home. This affects reliability, not replicability, as the study could easily be run again.

Question 3

- (a) The majority of responses could clearly outline one other aim of the study by Dement and Kleitman. The most popular choice was the dream recall differences between stages of sleep. A minority of responses repeated the aim already provided in the question. This is another example where candidates need to take their time and read the question carefully to ensure their response is focused on what the question is actually asking for.
- (b) The majority of responses could identify two other features of the sample. Popular choices included how many males participated and that only five were studied intensively. Common errors included presenting incorrect number of males and females (some responses presented sample sizes that did not add up to nine even though that number was in the question) or that some withdrew from the study.

- (c) The majority of responses could identify an appropriate application to everyday life. The most popular choice was helping people with sleep problems. However, many candidates did not then explain how this could be achieved. Stronger responses could explain that the EEG could be used to compare people with and without a sleep disorder to help diagnose. For these types of questions, candidates need to cover 'what' the application is and 'how' it can be achieved.

Question 4

Stronger responses could present a correct result with a meaningful comparison and correct data. However, this was only from a small minority of candidates. The majority of responses focused on target-present line-ups. This is another example of candidates needing to read the question carefully before presenting a response. The question was about target-**absent**, not present. Therefore, the majority of responses could not be awarded any credit due to this mis-reading. Candidates need to know all aspects of all 12 core studies to be able to focus on the demands of every question. This had the joint second-highest blank response rate on the paper.

Question 5

- (a) A majority of responses could clearly outline one assumption of the cognitive approach. The most popular choice was the computer analogy. There were some brief assumptions provided by some candidates with limited terminology linked to cognition. These could only be awarded partial credit. The assumptions for all four approaches are clearly outlined in the syllabus.
- (b) A minority of responses could clearly explain why the study by Baron-Cohen et al. is from the cognitive approach. These strong responses could provide a finding from the study and then give a clear explanation as to why it supported one of the assumptions provided in **Question 5(a)**. Many candidates wrote out the assumption from **Question 5(a)** again to 'explain' why Baron-Cohen et al. was from the cognitive approach and could not be awarded any credit as they had already been awarded marks in **Question 5(a)**. To improve, it is very important for candidates to read the entire set of questions from any question number, for example **5(a)** and **5(b)** to ensure that their responses are logical and follow the demands of the questions. This had the joint second-highest blank response rate on the paper.

Question 6

- (a) A minority of responses could clearly outline the term shaping with an example from the study by Fagen et al. Common responses included successive approximation and how the banana was used as a 'reward aid' to help shape the elephants' behaviours. Many responses focused on the basic mechanisms of positive reinforcement which was not answering the question about 'shaping' so could not be awarded any credit. Some responses could outline the term but did not provide an example from the study by Fagen et al or provided their own example which could not be credited.
- (b) Stronger responses showed clear understanding of the ethics involved in research using animals as participants. Tenzin was overwhelmingly the most popular choice of friend. Common ethical choices included minimising harm, numbers, and housing. Many responses could clearly outline why Fagen et al. could be considered ethical in nature with a range of different examples directly from the study. Popular examples were only using five elephants, that the elephants could walk away to minimise harm, and that rewards were used and not punishment. However, some responses used human guidelines and as a result could not be awarded any credit. Common errors in relation to this were informed consent and the right to withdraw. For the latter, the elephant walking away is not an example of the right to withdraw as this is an exclusively human ethical guideline. It is very important for candidates to know the difference between human and animal ethical guidelines. They are all listed in the syllabus.

Question 7

The average mark for this question was three. Stronger responses could clearly outline what participants were required to do in the doodling condition. These responses clearly showed knowledge of the procedure which is an essential skill for this examination. Some candidates were able to describe the procedure but not once mention what a participant in the doodling condition was expected to do. The procedure needs to primarily focus on what the participants actually did in the study. Some responses to this question focused on the materials given to the participants without any mention of that the participant was meant to do with them. Materials do not gain credit here.

Question 8

There were a range of responses to this question. Stronger responses could suggest many reasons why the people on the bus did not help from the newspaper story. Popular choices included if the person who collapsed was drunk, the race of the victim, the sex of the victim, and the number of other people on the bus who could help. Diffusion of responsibility was used quite effectively in stronger responses. However, some responses focused on why people **did** help which was not answering the question. This is another example where candidates need to read the question carefully to know what the 'background story' is and then respond appropriately. To improve on question types like these, candidate should be prepared to present examples from the study that link to the novel situation in the question.

Question 9

- (a) The majority of responses could describe around three features of the sample used in the study by Hassett *et al.* Popular choices included the sample size at various stages of the study, where they lived and the species of monkey. Some responses outlined generic ideas about the sample without specific details (e.g., the number who had hormonal treatment). Some candidates described the procedure rather than the sample of participants. It is essential for candidates to know the features of the samples used in all 12 core studies.
- (b) Stronger responses could clearly explain one similarity and one difference. Popular choices to compare the studies on included experimental designs, ethics, type of data collected, and primary research method used. To improve responses to this type of question, candidates need to choose comparison points that can be developed and explained, using examples from both studies to explain the similarity and/or difference. For example, explaining the experimental nature of both studies would involve explaining that cause and effect can happen in both studies with examples of controls from both studies to allow cause and effect be stronger. However, stating that each study had a different aim does not allow the response to be detailed so will always only achieve Level 1. Candidates need to choose carefully what the comparisons are ensuring that they are logical and can be explained fully, using examples from both studies. It is also very important to read the question to see what can or cannot be used in the response. In this case, the candidates were told not to refer to the sample, yet a number of candidates did use the sample in their responses and were therefore awarded Level 0.

Question 10

The strongest responses evaluated the study by Bandura *et al.* in depth and in terms of two strengths and two weaknesses with at least one of these points covering the named issue of generalisations. Common choices included ethics, generalisability, observations, reliability, validity, and quantitative data. These strong responses could explain why an element of the study was a strength or a weakness using specific examples from the study by Bandura *et al.* to explicitly support their point. These answers tended to score Level 5 marks. Candidates need to ensure that they follow the demands of the question, covering two strengths and two weaknesses, all in equal depth. Some responses did cover the four evaluation points but were brief or did not use the study by Bandura *et al.* as examples which meant the response scored in the lower bands. Other responses included three evaluation points that were thorough, logical, and well argued with a fourth point that was not in context which meant it could not be give Level 5. Candidates need to know that any description of the study does not gain credit in these type of questions as it is testing their evaluation skills *only*. There were some common errors evident in some responses, for example, that both the children and parents gave informed consent and that the children watched the model on a video screen. Both of these are incorrect for this Bandura *et al.* study. In addition, some responses are **still** following a GRAVE approach to this question (Generalisability, Reliability, Application, Validity, Ethics) when the 'A' part is not creditworthy for an AO3 question. A response that fails to have one evaluation point about the named issue can only score Level 3 (6 marks) maximum. There were many responses that briefly outlined strengths and weaknesses with only some being in context which is a Level 2 response. Any response that has no context cannot get above a Level 1 mark. In addition, many responses did use ethics in an evaluative sense but did not fully explain why it could be a strength and/or a weakness or tried to use animal guidelines. Some responses did not cover the named issue. To improve on this question, candidates need to plan carefully, choosing two strengths and two weaknesses with one of these being the named issue, avoiding real world application where possible. Each strength and weakness should be of equal length with an explanation as to why it is a strength or weakness with examples (plural) from the study to show clear understanding. An evaluation that is in depth tends to have at least two explicit examples from the named study for every evaluative point made. These are the requirements for a Level 5 response. The average response was Level 2 (4 marks) for this cohort.

PSYCHOLOGY

Paper 9990/22 Research Methods

Key messages

- Candidates need to ensure that their responses are focused on the questions within the exam paper. There was more than one instance where it was clear that candidates had misread the question and provided responses which were not creditworthy: for example, some discussed covert rather than overt observation for **Question 1**.
- Candidates need to ensure that they understand the expectations for different command words used on the paper; for example, describe and explain. For describe, candidates need to ensure that they provide a sufficient number of unique points related to the marks allocated to the question, whereas for explain candidates need to identify a particular feature/concept/theory and then **link** their detail point to the feature/concept/theory they have identified. Often, explain questions had poorer outcomes in terms of marks as a candidate just described.
- Candidates need to ensure that they are able to define/outline key terms within the specification such as 'situational variables' and know the difference between terms such as qualitative and quantitative data. This is also important in questions where research methodology is used. It was clear that in some questions which talked about methods such as longitudinal studies, or volunteer sampling, there was a lack of understanding about what the terms actually meant.
- Candidates need to ensure that they link their answers to the information given in the stem if asked to do so. Often candidates showed excellent understanding of named issues/studies but then lost marks for giving generic responses.
- Candidates seem to struggle with questions which ask them to make a judgement about the validity of something. Often candidates would just talk about it 'not measuring what it is supposed to measure' which is not enough and is a basic definition of validity. Candidates need to look in detail at the information given and make an appropriate judgement about validity **based on** that information not validity in general.
- It is worth noting that candidate responses for the 10-mark extended response question showed good knowledge and understanding of interviews. There were a number of thoughtful responses for this extended response question, and candidates should be commended for their performance.

General comments

This was the second March series for the new Psychology specification. Paper 22 scripts provided the full range of marks, showing a good level of knowledge and understanding across many areas of the specification. Where performance was limited, it was due to a lack of knowledge of key terms, or a misunderstanding of the demands of the question.

This was clear when looking at questions on situational variables (**7**), correlations (**9(c)(i)**), and measures of central tendency (**4(a)**). Candidate responses showed some gaps in knowledge and understanding of key studies such as when referring to the performance test in the Fagen et al. study, and the Perry et al. (personal space) study which meant that candidate responses failed to access the full range of marks. This was especially true of questions which required candidates to comment on the validity of a method or particular test. For future series, candidates need to ensure that they have a good understanding of command words, key research method terminology such as validity and reliability, and the studies which have been named on the specification.

Comments on specific questions

Section A

Question 1

- (a) This one-mark question required candidates to outline what is meant by an overt observation. To achieve the marks candidates needed to give a brief but accurate explanation of the term, such as 'participants being aware that they are being observed'. This was indeed the most common response seen.

Candidate performance on this question was pleasing, most candidates were able to achieve the mark available. When candidates did not achieve the mark, it was due to candidates saying the aim of the observation is known to the participants. This may not be the case, as only the presence/role of the observer is known; this was a small minority of candidates however.

- (b)(i) This two-mark question asked candidates to suggest one reason why an overt observation may be better than a covert observation. To achieve the two marks candidate responses needed to make a suggestion and then provide some detail about that suggestion. The important point to note is that for two marks there needed to be some form of comparison to covert observation. The most common way candidates achieved this was by suggesting that overt observation is **more** ethical as you are more able to obtain consent than in covert observation.

The vast majority of candidate responses were able to achieve at least one of the marks i.e., by saying that overt is **more** ethical (taken as a comparative) with many then going on to provide more detail for the second mark. Where candidates did not achieve the two marks it was either due to a lack of detail or not giving a comparative. Another point to note is that even in an overt observation, participants may be deceived, e.g., about the aim of the study, the key difference is that they know – i.e., are not deceived about *the role of the observer*.

- (ii) This two-mark question asked candidates to suggest **one** reason why an overt observation may **not** be better than a covert observation. Similar to part (i) above, candidates needed to make a suggestion and then provide some detail about that suggestion. As above, for two marks there needed to be some form of comparison to covert observation. The most common way candidates achieved the two marks was by suggesting that there was a higher chance of social desirability with overt than covert as the participant may respond to the observers' expectations.

Fewer candidates achieved the two marks for this question, mainly due to their misunderstanding about demand characteristics as a term. It is important that candidates **do not** talk about participants **showing** demand characteristics instead they need to talk about **responding** to demand characteristics.

Question 2

This three-mark question asked candidates to explain what makes this (unseen) investigation a longitudinal study. To achieve the three marks available candidates needed to make at least two unique points about longitudinal studies and then make at least one link to the study provided.

Candidate performance on this question was pleasing. Most candidates were able to get at least two marks from this question and many got the full three marks.

Those candidates that achieved the three marks did so by:

- Stating what a longitudinal study is (a research method completed over an extended period of time)
- Who they do the study with (the SAME participants used over that period)
- How does the study show these components (patients are going to see Dr Bryan to report their symptoms every week for a period of time).

Where candidates did not get the third mark it was due to the explanation not making it clear that the same participants were used over a period of time. This is a vital part of a longitudinal study therefore candidates would not usually get the three marks without this element (this could be stated in context as well which is fine).

Question 3

This two-mark question asked candidates to explain why the withdrawal of two participants (in Holzel et al.) could have reduced the validity of the study. To achieve the two marks candidate responses needed to explain why the withdrawal will have a negative effect on the study (i.e., mean a smaller sample size) and then provide some detail of why their explanation would have an effect on the validity of the study.

Candidate performance on this question was mixed. Those that achieved the two marks did so by suggesting that the withdrawal would lead to a biased sample/gender bias/smaller and then going on to relate this to validity such as the fact that there could be something specific about the candidates who withdrew which may have changed the composition of the results (such as personality/stress levels). However, whilst most candidates achieved the first mark quite easily, very few went on to achieve the second mark.

Question 4

(a) This two-mark question asked candidates to outline the **most** appropriate measure of central tendency to use for a maths test in which each question is equally difficult. The appropriate test for this question was a mean. If any other answer was given, then no marks could be accessed. To achieve the second mark the candidates needed to provide some link back to the question which links to their choice. This link mark could be gained for responses such as:

- Correct answers on the test
- Correct scores on the test
- Number of answers they got correct.
- Test scores
- Correct scores.

Unfortunately, there were far more errors than expected on the responses to this question. Many candidates' responses suggested wrong tests such as median, mode, or generic answers such as 'calculate the average'. Other responses would talk about histograms or standard deviation which are of course incorrect. This is unfortunate as this is a fairly straightforward question to achieve some marks and for future series candidates need to ensure they understand how the different tests are used within research methodology. Where candidate responses did achieve the correct response (mean) many would then go on to link that to the scenario which enabled them to achieve both the marks.

(b) This one-mark questions asked candidates to state the **most** appropriate type of graph to use for the given data. The definitive answer here was a bar chart/graph, and most candidates achieved this mark. Where candidates did not achieve the mark, it was inevitably because they suggested a histogram which is incorrect as the data given was not continuous but categorical.

Question 5

(a) This two-mark question asked candidates to suggest **one** way that the performance test (in Fagen et al.) could have lacked validity. Candidates needed to justify their answer. To achieve two marks candidate responses needed to:

- Make a suggestion as to why it may not be valid
- Justify their suggestion.

Unfortunately, candidate responses for this question were a little muddled and often did not relate to the question. This question required candidates to look at the test itself as a way of measuring performance rather than looking at why the elephants themselves did not perform well in the test. This was a common error with such responses about the elephants' personalities, health, and that the test was too easy/hard which were not creditworthy.

Responses which could gain credit included:

- The idea of only needing 80 per cent correct (which means we do not know what happened with the 20 per cent that was wrong and whether they were near misses or **very** wrong)
- The idea of subjectivity in observing the elephant meaning different observers could interpret the behaviour differently (lowering internal validity).

The second bullet point was the more common response but was still in the minority.

- (b) This one-mark question asked candidates to suggest **one** way that the validity of the performance test could have been improved. To achieve the two marks candidate responses needed to make an appropriate suggestion only.

Inevitably the strength of responses for part (b) was directly related to how the candidates had responded to part (a). If part (a) was incorrect it was very likely that part (b) was incorrect, which it was in the vast majority of cases.

Of those that did achieve the 1 mark it was usually due to them suggesting having the same independent observer for all elephants, or to ensure that they had to get a higher percentage correct to pass the test. There were very few correct responses seen, however, therefore increased knowledge about the Fagen et al. study is recommended for future examination series.

Question 6

This six-mark question required candidates to describe quantitative data and qualitative data, using any example(s). Awarded marks could come from accurate description of each term, further detail about the term and appropriate examples such as showing knowledge of the types of data of studies on the specification; such as the Piliavin et al. study having qualitative data due to the recording of the responses of the people in the carriage, or Milgram having quantitative data through the percentage of people who went up to 450 volts. Candidates performed exceptionally well on this question, with most having excellent knowledge and understanding of both terms and able to give appropriate examples whether real or made up by the candidates. The majority of candidates achieved most, if not all, the marks on this question.

At the lower end of the mark range, candidates usually achieved one or two marks for suggesting that quantitative data is numerical whilst qualitative data was in depth. They may be able to give an appropriate example such as from the Milgram et al. study. Candidate explanations of qualitative data were weaker and often were limited to suggesting that depth was gained through open questions. Even at this level many were getting 3 or 4 marks.

At the higher end of the mark range candidates produced some thoughtful responses and the vast majority were able to give detailed explanations of both terms highlighting a number of ways of achieving that type of data. The examples given were appropriate and explicitly linked to the question. These candidate responses often achieved full marks on this question.

Question 7

- (a) (i) This two-mark question, based on a novel scenario, asked candidates to suggest **two** situational variables that should be controlled so that they are the same for the two sections of the factory. To achieve the two marks on this question candidate responses needed to give two unique situational variables which could include: temperature, light, noise levels, type of work/job, how physical the job was etc.

Candidate performance on this question was average. Many candidate responses were able to achieve one mark but very few got two. Although some candidate responses unfortunately focused on participant variables, the most common error on this question was that responses focused far too much on the workstations/people at the workstation themselves. Although ensuring the structure of the workstation (meaning size of actual workstation itself) is the same is creditable, responses such as the people liking each other (participant) or the colour of the workstation (not really relevant in terms of a study on personal space) was not. Also, cleanliness of the workplace was not creditable as there is no real reason to suggest this would be different. Finally, although ensuring the number of people in each half is the same was creditable, the use of participants was not. The most common responses were temperature and light.

For future series candidates need to ensure that they know the difference between participant and situational variables, and that the response that they give is relevant to the question asked.

- (ii) This four-mark question asked candidates to explain for **each** of their suggestions in part (a)(i), why **each** situational variable would be relevant to Jacinda's study. For each of their variables candidate responses needed:

- **Explanation** of why it is relevant to the study
- And then provide further **elaboration** on this.

Inevitably, whether candidates achieved the marks on this question was largely dependent on whether they have got the first part correct. If candidates kept the first part simple and suggested such variables as temperature and light, then they were more likely to achieve most of the marks for part 7(a)(ii). Where the responses were more complicated and suggested such variables like the structure of workstations, what is on the workstation etc., then they have not really been able to explain the reason this is relevant (usually by stating how it affects results) for the second part.

The most common correct responses were:

- If the temp is hotter in one area then the participants will become more irritable (explanation) therefore they may want to step further away from others (detail)
- Increased light in one area may make them more irritable/may give them a headache (explanation) and therefore they may not want to be close to others (detail)
- Less light in one area may mean that they cannot see what they are doing (explanation) so may need to be closer to others to work together (detail).

These responses were in the minority, however.

- (b) This three-mark question asked candidates to suggest how Jacinda could obtain a volunteer sample for her study. For candidates to get the three marks available they needed:

- What they are using to get their sample to notice the study (so posters/flyers/email etc.)
- How will they achieve this in relation to her specific sample (so workplace noticeboard/email to work address etc.)
- How will she know they want to volunteer? (so, they reply to the email/she leaves contact details so they can reply).

Candidate performance on this question was fairly pleasing with most candidates able to achieve the majority of the marks. Where candidate performance was limited, it was usually due to the misinterpretation of the type of sampling needed (with some responses suggesting methods more akin to opportunity sampling), or not writing the need for contact details on a poster/that those that respond to an email would become the sample.

Despite this there were some very good responses which showed excellent knowledge about this specific type of sampling which was pleasing to see.

- (c) This two-mark question asked candidates to suggest how Jacinda could measure her participants' interpersonal distance preference, other than by using an interview. To achieve the two marks available candidates needed to:

- Suggest an **alternative** method to an interview i.e. questionnaire, rating scale, open question etc.
- Provide some **detail** linked to the scenario.

Candidates performed well on this question, with the vast majority able to give an appropriate method to use with many giving the linked detail to enable them to achieve the two marks. Where candidate performance was more limited it was due to a lack of detail for the second mark i.e., the response did not give a specific question or rating scale.

Note that observation was accepted as a response for this question but as it is quite difficult to observe interpersonal distance without any contact therefore responses struggled to achieve more than 1 mark. Occasionally a response would discuss categories such as expressions of disgust when people come close/physically moving their desk away from another etc., but that was the very rare. For future series it is worth emphasising that alternative methods need to be appropriate for the question asked, and some thought should be given to how feasible it really would be to perform a particular study using certain types of methodology.

Question 8

- (a) (i) This one-mark question asked candidates to suggest one ethical reason, other than housing, why it is important that the monkeys are comfortable in the environment in which they are housed when they are not being tested. There were no real issues with responses to this question with many using the ethical issue of pain and distress which was creditworthy. Physical harm was also creditworthy. Where candidates did not achieve the mark available it was invariably due to the use of human ethical issues such as protection from psychological harm, or just repeating the question itself. For future series is important that candidates understand the differences between human and animal ethical issues.

- (ii) This one-mark question asked candidates to suggest **one** methodological reason why it is important that the monkeys are comfortable in the environment in which they are tested. Creditworthy responses to this question included anything that suggested that the behaviour of the monkey may change due to lack of sleep/being frightened during the night etc. Candidate responses were pleasing for this question with the majority accessing the one mark available.

Where candidate responses did not achieve the mark, it was due to the vagueness of their responses which concentrated on terms such as reliability and ecological validity without explaining them in the context of the study. In addition, some responses suggested it was important, so they co-operate with the 'experimenter' and 'perform better', unfortunately this was not creditworthy as it is normal behaviour that is needed.

- (b) This two-mark question asked candidates to suggest **one** way to measure which material the monkeys prefer inside their cages. To achieve the two marks for this question candidates needed to suggest:

- A **way** we can measure preference (can be a method i.e., observe or a methodological detail)
- Further **elaboration** on the way suggested.

It was important that within the elaboration point some idea was given of how they are going to measure their preference. This did not have to be numerical; it could just be that they then use observation to see which one (of leaves or paper) they prefer to sleep on.

Other common responses seen were:

- Put cages side by side one with leaves and one with paper (1) See/observe which one the monkey sleeps on during the night (1)
- Use different materials in cages each night (1) then time in minutes how long they sleep in each one (1)
- Observe each night with one night paper and one-night leaves (1) see how many times the monkey wakes up during the night (1)
- Have a cage with both leaves on one side/paper on other (1) then see/observe which side they sleep on (1).

Candidate responses to this question were really pleasing with many candidates using variations of the above responses to achieve the two marks available. Where performance was more limited it was due to the vagueness of the responses. For example:

- Record how long the monkeys sleep on each material **or** number of hours spent interacting with the leaves or paper: – This only got one as they did not suggest what they are going to do with that material in order for them to be able to record this.

There were some really creative responses for this question, which meant that there were many full mark responses.

- (c) This two-mark question asked candidates to identify two features (other than the materials in the cage) that are important in the 'housing' guideline in relation to animal ethics. Candidates' performance in this question was pleasing with the vast majority of candidates achieving the two marks available. The most common responses for this question were space and food.

Question 9

- (a) This two-mark question asked candidates to explain why it is better to use a correlation than an experiment for this investigation. For the two marks candidates needed to:

- Give an **appropriate** explanation about why a correlation would be better **for this study**
- Provide some **detail** about this explanation in relation to **the study**.

The important point to note for this question is that it needed to be related to the study itself in some way. Although the first mark may be generic i.e., 'because the relationship may not be causal' or linked 'because it is unethical to manipulate a person's happiness' but it still needs to be appropriate for the study given.

Candidate performance on this question was below average. Although some were able to gain one mark for a fairly generic explanation of why a correlation would be better, few were able to provide the detail needed to achieve the second mark.

A response which could have achieved the two marks would be:

- Because it would be unethical to change a person's mood to make them happy/sad (explanation)
- Therefore, you would not be able to manipulate the IV of happiness which you would need to do in an experiment (detail).

Unfortunately, this was very rarely seen. For future series it is important that candidates note the possible advantages/disadvantages of using particular methods such as an experiment and correlation and have some practice at relating these evaluation points to novel scenarios and unseen studies.

- (b) This one-mark question asked candidates to suggest **one** way that exercise could be measured in this investigation. To achieve the mark candidates needed to suggest one way that would be appropriate for the study in the question. The type of responses which could achieve the mark would be: the number of steps, the duration of the exercise (by far the most common response).

Candidate performance on this question was mixed. Where candidates kept it simple, they were invariably able to achieve the mark. Where performance was limited, it was due to responses that typically did not actually give a measurement i.e. just saying observation or questionnaire without giving an actual measurement (such as observing how long a person exercised for, or a specific question they would ask).

For future series it is important that candidates are able to differentiate between when a question asks about a 'way to investigate' (in which responses such as observation would gain credit) and a 'way to measure' (in which an actual measurement is needed).

- (c) (i) This two-mark question asked candidates to suggest **one** strength of this way of measuring exercise. To achieve the two marks candidate responses needed to:

- Give an appropriate **strength** of the way suggested in part (b)
- Give some **detail** about the strength identified.

Candidate performance on this question was directly linked to the ways suggested in part (b). Where candidates had given an appropriate way of measuring exercise in part (b) they often were able to achieve both the marks in this part of the question. If they had not given an appropriate way/or not an actual measurement in part (b) invariably they were unable to gain the marks in this part. This was not always the case, however, as some responses which just suggested a method such as 'observation' in part (b) then went into further detail in this part to justify/elaborate on their strength which allowed the response to achieve the marks.

One point to note for future series is the strength of it being quantitative data was not creditworthy for this question. This is because in a correlation only quantitative data can be collected and measured anyway.

- (ii) This two-mark question asked candidates to suggest **one** weakness of this way of measuring exercise. To achieve the two marks candidate responses needed to:

- Give an appropriate weakness of the way suggested in part **(b)**
- Give some detail about the weakness identified.

Similar to part **(c)(i)**, candidate performance on this question was directly linked to the ways suggested in part **(b)**. Where candidates had given an appropriate way of measuring exercise in part **(b)** they often were able to achieve both the marks in this part of the question. If they had not given an appropriate way, or not an actual measurement in part **(b)** invariably they were unable to gain the marks in this part. As above, this was not always the case, however, as some responses which just suggested a method such as 'observation' in part **(b)** then went into further detail in this part to justify/elaborate on their strength which allowed the response to achieve the marks.

Question 10

- (a) This ten-mark essay question asked candidates to describe how Dr Anand could use structured interviews to study conversations that 8–10 year old children have with their parents. Candidate responses showed the full range of marks, with some really nice responses at the higher mark range. Candidates did still find it a challenge to achieve the higher levels (5) but there were a significant number of level 3/4 responses which was pleasing to see.

Many candidates had a sound understanding of interviews in general and were able to make relevant decisions about the type of questions used, and how to interpret the data they gained. Where candidates struggled slightly was with the format of the interview i.e. the structured element. Occasionally, there was slight misunderstanding of what was meant by structured interview with some responses suggesting that structured meant only having closed questions which is not always the case. One of the elements that candidates needed to achieve to get into the higher levels was to give examples of questions and candidates excelled on this part of their response, with many giving relevant and accurate examples of the questions they could use, and the type of data that this could achieve which was pleasing to see.

Candidate responses within the lower mark range were often able to give a list of questions which they were going to use and were able to produce a basic procedure to follow. However, at this level, candidate responses would have significant gaps within the procedure which would mean that it would not be replicable. Candidate responses may also mistakenly suggest other research methods that they could use. For example, some candidates mistakenly suggested independent and dependent variables and detailed observational or experimental procedures. In this case, the response could still achieve marks for suggesting types of questions, but it was inevitable that these responses struggled to get higher than level 1 or 2. At this level, the interpretation of data element was often missing or very basic i.e. a brief point about numerical scoring. At this mark range some candidates still talked about ethics and sampling which was not creditworthy.

Candidate responses within the higher mark ranges would be able to describe a procedure that would be replicable by other researchers. Most candidates at this level would suggest both open and closed questions, and then describe in detail examples of questions that could be used within the study. This would include answer choices for closed questions and the correct command words such as describe or explain for the open questions. Within this description candidates would highlight the type of data that would be produced by these types of questions. At this range, candidates would then go on to describe how their data could be interpreted such as the use of the mean/median, and bar graphs for the quantitative data and looking for themes within qualitative data. What differentiated candidate responses at this level was whether they explicitly talked about the technique used for the interview, understanding what was meant by structured interviews. Responses went further than just stating that it would be the same questions but also looking at other elements of their procedure such as keeping tone of voice, where the interview was done, timings etc. the same which was excellent to see. These candidate responses would be examples of those who would achieve high level 4, or even level 5. Very few candidates at this mark range discussed anything about ethics or sampling.

It is essential that candidates are prepared for this question and have a clear understanding of the four required features for each method they can be asked about. This will ensure that in future series candidates are able to achieve the marks at the highest levels.

- (b)(i)** This two-mark question asked candidates to explain how one part of the procedure you described in part **(a)** helps to make the study reliable. To achieve the two marks candidates needed to identify a part of their procedure that helps reliability and then provide further explanation for that point.

Candidate performance on this question was mixed. Where candidates achieved the two marks, they identified appropriate parts of their procedure such as using the same questions with everyone and the same procedure and were able to give relevant detail for the second point (such as the fact that the same questions mean it is replicable to test reliability).

The most common error in responses was a lack of detail meaning that often the second mark was not available. Unfortunately, when candidates had completed experiments/observations as part of their procedure, they often would discuss these in this part of the question but inevitably this would not be creditworthy.

One point to note, the use of inter-rater reliability is still causing some issues. Even if you have two interviewers/people to interpret the data that on its own does not mean you have inter-rater reliability. This will only occur if the results gained by the two observers are the same/very similar. (above 0.80 in terms of correlation). This is a common error and needs to be addressed for future series.

- (ii)** This two-mark question asked candidates to explain how one part of the procedure you described in part **(a)** could be a problem for the reliability of the study. Similar to part **9(b)(i)**, candidates needed to identify a part of their procedure that may cause a problem for reliability and then explain/exemplify with this point. Again, candidate performance on this question was mixed, and the clarity of responses was directly related to the detail they had given in their procedure

A number of candidates produced thoughtful responses for this question, with issues such as the use of open questions (where appropriate) and then go on to say how this made the data gained through the interview difficult to replicate/interpret meaning that reliability is lowered, or that there may be extraneous variables affecting reliability if they did not keep certain elements of the procedure standardised.

PSYCHOLOGY

Paper 9990/32 Specialist Options: Approaches, Issues and Debates

Key messages

Questions 1, 3(a), 5, 7(a), 9(a), 9(b), 11(a), 13(a), 13(b), 15(a) –

These questions in this exam asked candidates to apply an area of the syllabus (theory, technique/treatment, self-report, etc.) to explain how it is relevant to a particular scenario or context. It is important that candidates are aware of the titles of the bullet points in the syllabus. It would be helpful for candidates to do revision notes with the title of the topic area and bullet point at the top so that they can identify which part of the syllabus these types of questions are referring to. Candidates should also refer directly to the scenario/context in the question in their response.

Questions 3(b), 7(b), 11(b) and 15(b) –

These questions in this exam asked candidates to evaluate the suggestion such as the technique/treatment that was outlined in the candidate's response to part (a) of the question. In this exam, these types of questions asked the candidate to evaluate the technique outlined in part (a) such as with a weakness, explain a practical application of part (a) or a problem with the technique outlined in part (a). It would be helpful to candidates when doing revision to learn strengths and weaknesses of the theories, techniques, self-reports, treatments, etc., they have learned and put these into their revision notes. They should also practice explaining the evaluation point in the context of the question.

Questions 2, 6, 10 and 14 –

Part (a) – These questions could ask the candidate to outline a theory, study, technique/treatment or self-report used by psychologists that is named in the syllabus or outline one of the issues and debates, possibly with an example from the syllabus content. The revision technique outlined previously in this report will aid candidates to learn the syllabus material.

Part (b) – This part of the question may ask candidates to explain a strength or a weakness of the issue/debate or the syllabus content outlined in **part (a)**. The question could also ask candidates to explain how a bullet point in the syllabus links to or supports one of the issues or debates. It would also be useful for candidates to write revision notes where they define the issues/debates and prepare a strength and a weakness of each issue and debate to prepare for the **part (b)** of this type of questions. Candidates should also note how the topics covered in the syllabus fit with each of the issues/debates. These questions in this exam were worth 2 marks for each part of the response and therefore a short response is appropriate.

Questions 4(a), 8(a), 12(a) and 16(a)

These questions in this exam came from one or two of the bullet points in the syllabus. This exam either asked the candidate to outline a key study from the syllabus or two studies, theories, characteristics/explanations/treatments of disorders or techniques identified in the specification under the appropriate bullet point. For this exam, some of the answers used the incorrect topic area in the syllabus or the description was brief. It could be useful for candidates to create revision notes with the title of each topic area and the description in the bullet point as the header. Alternatively, candidates could create a mind map and put this information in the centre.

Questions 4(b), 8(b), 12(b) and 16(b)

This question will always ask the candidate to evaluate the studies, theories, characteristics/explanations/treatments of disorders or techniques described in **part (a)** of the question. The response must include at least two evaluation issues, including the named issue, in order to be considered to have presented a range of issues to achieve the top band. However, most responses that evaluated using two issues in this exam, achieved in the lower bands due to the response being superficial and often with little analysis. Some responses that considered three issues tended to achieve higher marks as these responses were able to demonstrate comprehensive understanding with good supporting examples from the studies, theories, characteristics/explanations/treatments of disorders or techniques described in the **part (a)** of the answer.

The candidate must also provide some form of analysis to access level 3 and above. This could be done by discussing the strengths and weaknesses of the issue being considered, presenting a counter-argument to the issue under discussion or comparing the issue between two studies and/or theories. The response needs to explain the comparison/strength/weakness or counter-argument with examples from **part (a)** of the question. It was common for responses to state that two theories, for example, were 'similar' or 'in contrast' for an issue without any explanation as to why they could be compared in this way. This is limited analysis. A conclusion at the end of each issue would be helpful in order to show excellent understanding of the issue under discussion. In order to achieve the requirements of the level 4 and 5 descriptors, it would be best to structure the response by issue rather than by study and/or theory. It would also be ideal for the response to start with the named issue to make sure the answer covers this requirement of the question.

A small minority of candidates did not evaluate using the named issue. Quite a few of the answers were structured by study/theory/treatment rather than by the issue which often led the response to be quite superficial and repetitive. A number of the responses did include analysis. Candidates should be aware this question is worth 10 marks and need to include an appropriate amount of information.

General comments

The marks achieved by candidates for this session of the 9990 specification achieved across the full range of the mark band which was very pleasing to see. Some candidates were well prepared for the exam and showed good knowledge, understanding, application and evaluation throughout their responses. A significant number of candidates were not as well prepared and showed limited knowledge and understanding with brief, superficial and sometimes anecdotal responses. These candidates often had limited evaluation and application skills.

Time management for this paper was good for the majority candidates and most attempted all questions that were required. A number of candidates did not respond to one or more of the questions asked in the option area. A very small number of the candidates attempted to respond to more than two topic areas but often did not attempt all of the questions for each option chosen. These responses achieved at the lower end of the mark band.

The questions on clinical were the more popular choice of option, followed by organisations.

Comments on specific questions

Clinical Psychology

Question 1

The responses to this question covered the full range of marks. Responses that achieved 3–4 marks often referred to the process of CBT and how this could help to treat Rahul's symptoms specifically with accurate terminology. Common suggestions for contextualised responses involved Rahul would identify his irrational thoughts/obsession(s) and how this was leading to the compulsive behaviour of arranging things in order. He would be given homework to try to not arrange things in order and to see that nothing bad happened to his family. He would then experience fewer of these thoughts which would lead to a reduction in the compulsive behaviour. Responses that outlined exposure and response prevention were also able to achieve full marks if clearly contextualised to Rahul. Weaker responses lacked detail and often did not engage with the stem about how CBT will treat Rahul's OCD.

Question 2

- (a) There were some good responses that were able to clearly outline what is meant by determinism, including an example from the psychodynamic explanation of fear-related disorders to achieve full marks. Weaker responses were often able to give a definition of determinism but could not explain how the psychodynamic explanation of fear-related disorders was deterministic. Responses that did not receive credit included those that outlined what is meant by reductionism rather than determinism. Some gave an example from behaviourism and the 'little Albert study' which was also not creditworthy.
- (b) There were some clear explanations given of one problem psychologists may have when investigating the psychodynamic explanation of fear-related disorders. Common full mark responses explained the difficulties of investigating a non-observable phenomenon or that the data collected is usually qualitative (through conversations with the therapist/psychologist) and is subjective. These often had an example from the psychodynamic explanation such as that the unconscious, ID, ego, super-ego cannot be measured in a scientific way which helped the response to achieve full marks. Weaker responses often gave a brief answer which just identified the problem without clearly explaining it. Those that wrote about the incorrect explanation in **part (a)** often did not achieve credit in this part of their answer.

Question 3

- (a) Responses to this question covered the full range of the mark scheme. Those that achieved 3–4 marks often were able to describe the scoring on the GAD-7 and were able to explain how the test could be applied before therapy (to get a baseline), during and then afterwards. Weaker responses often just stated when the GAD-7 could be used (before and after therapy) and that it measured the severity of anxiety with no more details given. Some responses confused the GAD-7 with other measures of assessing anxiety or gave incorrect scoring which was not creditworthy.
- (b) There were some full mark responses to this question that could relate the weakness of the GAD-7 to either features of this self-report specifically or specific features of generalised anxiety disorder. However, most responses identified a weakness (such as social desirability or lack of detail due to it being quantitative data) but did not link this to either the GAD-7 or a feature of generalised anxiety disorder so achieved 1 mark.

Question 4

- (a) Responses varied for this question and covered the full range of the marks available. Level 3 responses were able to outline both the biochemical (dopamine hypothesis) explanation and the psychological (cognitive) explanation of schizophrenia. Level 3 responses were able to give specific areas of the brain and explain the dopamine levels and link these to specific symptoms. For the cognitive explanation there was a clear outline of failure to self-monitor and lack of theory of mind which was also linked to the symptoms. Weaker responses gave either brief details about dopamine levels, sometimes with inaccurate areas of the brain, or the symptoms in terms of positive and negative symptoms and the relationship to dopamine. For the cognitive explanation these often just outlined the failure to self-monitor/not recognising the hallucination as their own internal voice with few other details. Some attempted to outline a study for one or both explanations, sometimes with good accuracy and for the weaker responses these often lacked detail.
- (b) The marks for this part of the question did cover the full range of marks available with the most frequent levels awarded level 2 and 3. Those that achieved level 3 and above structured their response issue by issue and often started with the named issue of reductionism versus holism, along with evidence from the explanations outlined in **part (a)** and analysis. Apart from the named issues, other popular issues covered were determinism versus free will, idiographic and nomothetic and application to everyday life. Popular examples for the named issue included how biochemical explanation was reductionist and why whereas the cognitive explanation was more holistic as it does consider the biological cause but also looks at the cognitive issues patients with schizophrenia have.

Weaker responses achieving level 1 or level 2 did not contextualise their response. While some of the evaluation points were valid, the lack of context stopped responses from achieving a higher band. Some who wrote about a debate such as determinism versus free-will did not explain how

the explanation supported the relevant side of the debate. Some provided too many issues with no depth in explaining why.

Consumer Psychology

Question 5

There were several good responses to this question and some achieved full marks. Full mark responses were able to outline how the noise from the kitchen may affect perception of food by the customers. These gave a detailed answer with clear understanding of how noise from the kitchen may affect perception of food by the customers. The most popular responses included that the food would be more crunchy, less sweet and less salty. Some did refer to the results of the study by Woods et al. and link this to the experience the customers would have in the restaurant scenario outlined in the stem. Weaker responses often just outlined that the food would be less tasty/intense to achieve 1 mark. Some were confused and said that salty foods would be saltier or that crunchy foods would appear softer which was not creditworthy.

Question 6

- (a) Good responses were able to outline one explanation for why product placement in films affects choice. Auty and Lewis was often used to answer the question and this was effective when used to explain why the results occurred. Another common response that often did achieve full marks was that repeated / mere exposure to a product leads to increased positive feelings towards it and this led to purchase of the product. A large proportion of responses referred to a social learning theory by saying that the product was associated with a specific high-profile individual, which was not creditworthy.
- (b) Well answered by many. Full mark responses referred to psychological harm in some detail. Most common responses were that children are vulnerable and need to be protected. Weaker responses gave brief answers and commonly just stated that children needed to be protected from harm. A very few candidates were able to develop this point though most referred to protecting children from harm. Some responses did state consent. However, there was no credit for 'cannot get consent from children' without stating that consent can be obtained from a parent/carer.

Question 7

- (a) There were a number of full mark responses with two clear suggestions of ways that Malika could use Lauterborn's 4 Cs marketing mix model to attract younger customers. There were some strong responses, where candidates explained well why certain tactics would work for younger customers e.g., use of social media, which young people use commonly and linked these to convenience or communication. Some were able to explain cost with a specific link to younger customers but many just identified that she should lower the cost without any link which was not creditworthy. Weaker responses often gave a very brief suggestion and sometimes mixed up the 4Cs with the 4Ps which led to lower or no credit. Some misread the question and answered as if it were for customers over 50, which was not creditworthy.
- (b) There were a few full mark responses to this question which were able to explain one weakness of Lauterborn's 4 Cs marketing mix model. The most common creditworthy issue discussed was cross-cultural application and reductionism with an example from the model. A very common response was that the model lacks temporal validity because it fails to consider modern technology but as the response to **7(a)** often showed the model can easily be used in the modern context with reference to advertising on the internet. This issue was not creditworthy.

Question 8

- (a) There were many good, level 3 responses to this question with many showing very good knowledge of the Robson et al. study. Lots of responses gave good details of the aim, sample, questionnaire detail, spacing of tables (6, 12, 24 inches or 15, 30 or 50 cm), purpose of meal (romantic, business, friend), at least one result with many giving several results and a conclusion. Weaker responses gave fewer details of the study with some giving the incorrect sample or table spacing to achieve fewer marks. There were a significant minority of responses that outlined the study as a field experiment rather than a web-based questionnaire and these types of responses were often able to achieve very few marks due to this lack of knowledge of the study.

- (b) There were some level 3 and above responses to this question. Most responses evaluated using the named issue of self-reports with a discussion about the strength(s) and weakness(es) with good examples from the study. Other evaluation points included generalisability/cultural differences, ecological validity, qualitative and quantitative data and application to everyday life.

Weaker responses lacked depth in their discussion of self-reports with some responses stating the study had qualitative data which was incorrect. There was some confusion in weaker responses about the Robson et al. study with many incorrectly stating that it took place in a restaurant (when it took place as an online survey) so no marks were awarded for the discussion around good ecological validity. Weaker responses also provide a long list of evaluation issues with just a vague connection to the study outlined in **part (a)**.

Health Psychology

Question 9

- (a) Many responses achieved full marks by giving two correct physical effects of stress that Samay may be experiencing. Common responses included fatigue, weakened immune system, headaches and high blood pressure. Weaker responses gave incorrect physical effects and some identified psychological effects which was not creditworthy. Some candidates mentioned stomach ulcers. This is a persistent myth (stomach ulcers are caused by helicobacter pylori bacteria and long-term use of NSAIDs) and this symptom was not creditworthy.
- (b) There were many strong responses to this question achieving full marks. Better responses linked the suggestion to one of the physical effects identified in **part (a)** and linked this directly to issues Samay may have at work. For example, a common physical effect chosen was fatigue. Responses outlined that he would find it difficult to arrive at work on time and/or difficult to concentrate while at work due to sleepiness/lack of energy. There were a few responses that were weaker often because they were just very brief such as difficult to contrate due to fatigue. If no marks were awarded for **part (a)** it was rare for marks to be achieved in this part of the question.

Question 10

- (a) Good responses were able to define the idiographic approach and apply it to non-adherence well. The best of these referred to the Laba et al. study. Weaker responses often lacked the link to/example of non-adherence to medical advice for the second mark. Some outlined the nomothetic approach which was not creditworthy.
- (b) There were some good responses to this question with a clear explanation of one weakness of taking an idiographic approach to understanding non-adherence to medical advice. The most common weakness was that an idiographic approach is unable to produce general laws/predictions about human behaviour and the effect this could have on understanding non-adherence such as being unable to have a general approach to resolving the issues a patient has with non-adherence. Weaker responses tended to be brief and often did not link to non-adherence. Lack of understanding of the idiographic approach and/or non-adherence were quite common leading to no marks being awarded.

Question 11

- (a) A small number of responses were able to achieve 3–4 marks by giving a detailed suggestion about how Dr Singh could use contracts to improve the adherence of his patients to regular exercise. The strongest responses suggested what could be included in the contract and how it should be signed by the patient. Weaker responses lacked detail and often explained why contracts are effective which was not answering the question and was not creditworthy. In addition, several responses suggested the use of a reward, which would not be practical or remotely realistic such as giving the patient £100 for completing the contract. These type of responses were not creditworthy.
- (b) There were some full mark responses to this question that explained a practical problem with the suggestion given in **part (a)**. The most popular response was stating that the patient may say they are attending the class when they are not. Weaker responses were often not clearly linked to contracts (for example, stating the patient may lie) which tended to achieve 1 mark. Some responses stated that the patient would not sign the contract which does not answer the question

which was asking about a problem with implementing it. If the patient did not sign the contract then it would never be implemented. This type of response was not creditworthy.

Question 12

- (a) The responses to this question covered the full range of the mark scheme. Better responses gave clear and often detailed and accurate description of a study using fear arousal to improve health, and a study about providing information so people know how to improve their health. Popular studies used were Janis and Feshbach for fear arousal and either Tapper et al. or Lewin et al. for providing information. Strong responses were able to give an outline of the procedure of the study as well as a result. Weaker responses often gave fewer or some incorrect details. A common error was to state that the high fear arousal condition had the most change in behaviour in Janis and Feshbach's study which was incorrect. A few responses outlined a study on positive psychology which was not creditworthy.
- (b) The marks for this question covered the full range of marks although many achieved Level 2 and below due to lack of contextualising their discussion and/or lack of analysis. Most responses attempted the named issue of longitudinal studies and some could give a strength and a weakness and used the studies outlined in **part (a)** as examples. Weaker responses often did not refer to a specific study and would commonly state that longitudinal studies take a long time with no other points raised. Other common issues included self-report, ethics, reliability, validity, quantitative data and application to everyday life.

There were a number of weak responses to this question with the responses giving very brief strengths and weaknesses for the studies from **part (a)** without any clear examples from the studies to back up their points or any analysis.

Organisational Psychology

Question 13

- (a) Full mark responses often focussed on how Jaya's staff could be passionate about books and a love of reading, demonstrating high job involvement. Weaker responses gave a basic suggestion and did not link this to the scenario outlined in the question of being in a bookshop. These were often to state that those with high job involvement preferred to work alone. Some responses suggested it was connected to interactions with other staff members or Jaya as being a factor in their job, which does not address the question and was not creditworthy.
- (b) Common responses to this question included not agreeing with the goals of the bookshop and those that achieved full marks linked this to the scenario. Another common response was to suggest how alternative employment opportunities would attract staff away such as for higher pay. Those that achieved full marks did put this into the context of the bookshop. Some responses gave a vague suggestion about how the staff interact with other staff and/or Jaya with no mention of preferring to work alone rather than in groups. These types of responses were not creditworthy.

Question 14

- (a) There were some good, full mark responses to this question. Those that did were able to give two accurate features of the job descriptive index (JDI). Most common responses included the JDI having 72 items, the scoring system used and measuring satisfaction with pay, promotion opportunities and relationship with co-workers. Weaker responses often identified pay as one feature but gave no further features or incorrect features. Some responses gave incorrect details of the number of items and/or the scoring system which were not given credit.
- (b) Better responses were able to explain one practical application of the job descriptive index (JDI). The most common response explained how employers could give the JDI to their employees, find areas which were unsatisfying (e.g. opportunities for promotion/pay) and make changes to their company. Weaker responses often stated the organisation could identify areas for change but did not explain what they would do once the area had been identified. There were a significant number who suggested that the JDI could be given to prospective employees to help the organisation find the best workers. This is not a creditworthy application of the JDI.

Question 15

- (a) There were a number of good responses to this question with many giving two clear suggestions of how the manager could use two of Thomas-Kilmann's conflict-handling modes to reduce the conflict about lunch breaks. Common responses that achieved 2 marks per suggestion were about compromise, collaboration or accommodation and gave specific examples of reaching a solution regarding the lunchtime routine. Several responses suggested that the production line in the factory should simply stop for an hour, which would clearly be impractical so was not creditworthy. Others suggested that some employees would not have a lunch break at all which was also not a creditworthy suggestion.
- (b) Many responses were able to offer a weakness for one of the modes to reduce conflict suggested in **part (a)**. Stronger full mark responses effectively discussed a weakness of the technique leading to more conflict in the future or employees who have not got their way regarding the lunchtime will be more dissatisfied and an issue this could lead to on the production line. Weaker responses often just identified that further conflict might occur without linking it to the mode suggested in **part (a)**. Those who did not achieve marks in **part (a)** were unlikely to achieve any marks for this part of the question.

Question 16

- (a) There was a range of responses to this question covering the full range of the mark bands. Many responses gave clear and detailed descriptions of both universalist theories of leadership, and Heifetz's six principles in meeting adaptive challenges. Some of the weaker responses were able to outline Great Man Theory but gave vague or no details of charismatic and/or transformational leaders. Weaker responses sometimes just listed the six principles for Heifetz or gave an outline of one or two principles.
- (b) The marks for this question tended to be between level 1 and level 3. Many responses were able to correctly identify that Great Man Theory supports the nature side of the debate and many could explain why, with some contextualised examples, Heifetz's six principles supported nurture. Other common issues raised included determinism versus free will, reductionism versus holism, idiographic and nomothetic and application to everyday life.

Weaker responses often simply stated strengths and weaknesses of the theories, they rarely yielded high marks as points were not developed. Candidates should be encouraged to focus on a few issues and develop these well rather than simply stating that, for example, universalist theories are nature while Heifetz's six principles are nurture and moving on to the next issue.

PSYCHOLOGY

Paper 9990/42 Specialist Options: Application and Research Methods

Key messages

- What has been learned from the AS component of the syllabus should be transferred to the A2 component. For example, at AS candidates learn about methodology, such as experiments, which also apply to A2.
- Questions should be read carefully ensuring that the focus is on what the question asks.
- For **Section A** answers, candidates should relate their answer to the study in question or include an example. Questions frequently end with 'in this study' and so the answer should be related to that specific topic area/study.
- All terminology should be explained. Writing 'it is valid and reliable' for example, is insufficient without explanation, application or example.
- The syllabus includes 'example studies' such as 'e.g., Oldham and Brass (1979)'. Example studies can be substituted for alternatives, but these alternatives must cover the same or very similar content to the example study. If the Oldham and Brass study is substituted, the alternative study must be about a move to open plan offices and the data that was gathered from that move. The alternative cannot be about something different.

General comments

Some candidates answered questions from one option only. Other candidates, who correctly answered two options, sometimes performed considerably better in one option than the other.

Many candidates answered two questions from **Section B** instead of one (only one of these **Section B** responses can receive credit). Candidates are advised to read the instructions on the front cover of the question paper and to read the heading instructions for each question section.

Candidates should double check that the terminology they use in their answers is correct. Often terms such as reliability and validity were used interchangeably, as were qualitative and quantitative, and independent and dependent variables. There was also confusion with the terms format and technique in relation to questionnaires and interviews.

Section A

Question **part (c)** requires a general evaluative point that could relate to any study (such as a strength or weakness of a method) but it also requires for the general point to be related to the specific sub-topic/study in the question. Answers often included strengths and weaknesses, but these were often not related to the question, and this restricted marks.

Candidates should not use psychological terms without explanation. Frequently answers were limited to 'it is reductionist' or 'it is useful in everyday life' without further explanation. Stating 'it is reductionist' does not automatically identify it as a strength or weakness.

Candidates should not use the terms reliability and validity to answer every **part (c)** question for three reasons: (i) they do not apply to most questions and so cannot be awarded marks, (ii) candidates using the terms often do not know how they apply to the specific question and (iii) candidates often confuse the terms.

Section B

Candidates should only answer one question from this section.

Many candidates appeared to assume that they must conduct an experiment whatever the question. An interview, questionnaire or observation are methods independent of an experiment and candidates should not try to make other methods 'fit' into an experimental format.

Some candidates evaluate their plan in **part (a)** by listing strengths and weaknesses. This should not be done because: the question does not ask for evaluation; there are no AO3 marks allocated to evaluation; evaluation is done in **Questions (c)(i), (c)(ii) and (c)(iii)**.

Some candidates included a paragraph of results. This achieves no marks because the question asks for a plan only. Further, the proposed plan has not been carried out, so no actual results are gathered.

Candidates need to know the distinction between questionnaire format and technique, and interview format and technique, as stated on the syllabus: Questionnaire technique: paper and pencil (i.e., done by a person with the researcher present), online or postal. Questionnaire format: open and/or closed questions. Interview technique: telephone or face-to-face. Interview format: structured, semi-structured, unstructured.

When using psychometric tests candidates should not use acronyms unless the full title of it is provided first. For example, 'Beck Depression Inventory (BDI)' is fine, with BDI used afterwards. Further, it is insufficient to simply state 'I would use a questionnaire similar to K-SAS' (such as when writing about pyromania, for example).

Answers to **part (a)** questions in this section should include an appropriate plan, have applied a range (four or five) of specific (to the named method) methodological features, each of which should be explained fully, to show good understanding. Candidates should also include appropriate 'general' methodological features such as sample, sampling technique and location of the study. Many answers listed features such as 'I would have a random sample' and 'It would be an independent measures design' without explanation of why it would be a random sample, or how it would be obtained. Elaboration of these general sentences should be included.

In **part (b)(i)**, candidates should describe some relevant psychological knowledge that the whole question is based on. If the question, for example, asks about ways in which pain can be measured, then candidates should describe relevant measures.

In **part (b)(ii)**, candidates should explain what aspects of this psychological knowledge their **part (a)** plan is based on. These two question parts must be linked.

In **part (c)**, candidates must refer to what they did in their specific plan rather than give a generic answer that could apply to any study. Use of an example or quoting from their plan would be ideal.

Section B can be considered as follows: A teacher teaches a sub-topic from the syllabus and gives the candidate some psychological knowledge. The teacher then tells each candidate to plan a study using method 'x' to investigate some part of that sub-topic. The candidate plans the study using the psychological knowledge of the sub-topic and they use their methodological knowledge about method 'x'. In the examination, **part (a)** is the plan; **part (b)(i)** is the sub-topic knowledge and **part (b)(ii)** is how the knowledge was used to construct the plan. Exam question parts **(c)(i), (ii) and (iii)** then ask about some methodological decisions and evaluation about the plan.

Comments on specific questions

Section A

Question 1

- (a) Many candidates could be awarded full marks because they identified four features of the sample of participants. Correct answers (amongst others) included: 42 in the test group and 40 in the control; age range from 31–70; all from Hospitals in Croatia; all were diagnosed with bipolar type 1 or not.

- (b) Nearly all candidates were awarded full marks. Many stated that the reason why a control group was used was so that there could be a comparison between the test and control group (which earned 1 mark) and most candidates went on to state that this was so the alleles/polymorphisms/5-HTT2c and 5-HTT genes could be checked, any of these showing appropriate knowledge of the study.
- (c) A few candidates misunderstood the question and wrote about the sampling technique rather than the sample of participants. Candidates addressing the sample of participants were often awarded marks for including the fact that all participants had been clinically diagnosed by psychiatrists; that all participants were age and sex matched; or that there was a wide age range. Candidates linking each strength to the study were awarded the additional marks, for example by following up the 'wide range' strength with the comment that all participants were aged between 31 and 70 years old.

Question 2

- (a) Some candidates knew the ICD-11 criteria for gambling disorder and stated two features, for which full marks were often awarded. Some candidates could not be awarded marks because (i) they wrote about the general features of impulse control disorders or (ii) wrote about the general features of addiction. Both these answers could have been answering questions on kleptomania or pyromania with nothing specific to gambling disorder.
- (b) The question stated 'other than by self-report' yet many candidates decided to apply questionnaires and interviews, both of which are self-reports, and so no marks could be awarded. Answers scoring partial marks were often too vague, writing for example 'I would conduct an observation to observe gambling behaviour' without elaboration. Candidates awarded full marks stated a feature of an observation, such as 'covert observation' and stated what gambling behaviour would be observed, such as 'the number of times per week the person gambles'.
- (c) Many candidates provided generic statements such as; the gambler would not give honest answers; which were not incorrect, but which needed further explanation for full marks to be awarded. Some answers were more basic than this, such as; people might not give honest answers to questions; which didn't even use the words 'gambling disorder' and could apply to any question. Candidates must relate the answer to the question, in this instance gambling disorder, for full marks to be awarded.

Question 3

- (a) (i) A few candidates wrote about ethics in general, such as informed consent and debriefing, and by ignoring the question on deception, could not be awarded marks. On the other hand, there were many answers which identified two ways participants were deceived showing good understanding of the Hall et al. study. These included 'presenting themselves as independent consultants' rather than researchers conducting a study, and 'inviting participants to take part in a quality control test of jam and tea'.
- (ii) Candidates answering **3(a)(i)** correctly nearly always continued to score full marks here by giving two reasons for the deceptions. These included that it was necessary because participants would realise that they were participating in a study; because participants would realise that they were experimenters conducting a study; because the study would simply not work if participants were not deceived about the magic card trick.
- (b) Although there were some partial answers, with nothing more than 'it would not be valid' many candidates went further and stated the reason why the study would not be valid such as 'because participants would realise that they were being tricked by the jars or card trick and not behave as they normally would'. This latter answer is clearly linked to the study and would be worthy of full marks.
- (c) Like other **part (c)** questions, candidates often gave a strength and weakness without relating it to the question. Answers such as 'a field experiment has ecological validity' and 'variables are more difficult to control' were common and without any reference to the study these answers could not be awarded more than partial marks. Those linking their strength and weakness to the study, were often awarded full marks.

Question 4

- (a) All questions require some psychological knowledge to be shown for marks to be awarded, and so candidates guessing or stating how they would wrap a gift were nearly always awarded 0 marks. The syllabus states 'types of wrapping' and the two main types of wrapping this includes are 'traditional', where the gift meets expectations of looking like a gift with ribbons, bows, etc and 'non-traditional' where it may be difficult to determine that the gift is actually a gift, for example, through the use of plain brown paper wrapping. Many candidates were awarded full marks for outlining these two types.
- (b) By the answers given by candidates, most did not know what the term 'unstructured observation' meant, despite this being an essential feature of observations as listed on the AS syllabus. Candidates instead wrote about naturalistic or overt/covert observations or just opted to outline how an 'observation' could be used. It is advised that this is a *research methods* paper and so candidates should have full knowledge of every component of every method listed on the syllabus.
- (c) The lack of knowledge about the features of observations evident in (a) was also evident here. The syllabus states 'describe *the main features* of an observation (e.g., overt/covert, participant/non-participant, structured/unstructured, naturalistic/controlled). For this question candidates could use two strengths of any of these eight features to relate to gift wrapping. For example stating 'a covert observation means the participant is unaware they are being observed and wraps the gift as they usually would', or 'using a structured observation means that observers know exactly what they are looking for to see what features are included in a traditional wrapping, such as a bow'.

Question 5

- (a) (i) Many candidates confused results, findings and conclusions and answers which did not include conclusions could not be awarded marks. Results could be used to support a conclusion and often their inclusion enabled candidates to be awarded a second mark. The main *conclusion* was that 'emergency department physicians significantly underestimate pain' (awarded 1 mark) and for a second mark candidates could add 'from all medical conditions in paediatric patients ≥ 3 years old (+1 mark), especially from wounds, infections and soft tissue injuries, but less from fractures (+1 mark). What was not credited, for example, was the *result* that 'physicians assessed the child's mean pain to be NRS=3.2 (SD 2.0), parents: NRS=4.8 (SD 2.2) and children: NRS=5.5 (SD 2.4).
- (b) Many candidates suggested using the UAB pain scale, which could only be credited if appropriate parts of it were isolated (because it is a rating scale), such as observation of non-verbal cues and examples of grimaces and limping. Observation of verbal cues could also be credited, such as groans and crying. A clinical interview was also creditworthy where a medical practitioner could ask questions to achieve a subjective estimate of a child's pain.
- (c) There were many general answers that did not link to or relate to the question, as is common in all **Section A part (c)** questions. In this instance a candidate might write 'children may not understand complex instructions' which is not incorrect, but in relation to what? What study? By adding a simple comment such as 'and not be able to describe their pain accurately because they are in pain' would be sufficient for 2 marks to be awarded. All answers should always include the *in this study* component of the question.

Question 6

- (a) Most candidates were awarded full marks for their answers to this question. Malingering is when a person deliberately makes themselves unwell, or when a person pretends to be unwell (fakes an illness). However, many candidates stated that this was to 'avoid prison', a 'non-clinical' reason, rather than focus on malingering and Munchausen syndrome which is a mental disorder.
- (b) This question, like all other **part (b)** questions, invited candidates to think and suggest for themselves, rather than recall knowledge. This appeared to confuse some candidates. Three possible answers are: (i) conducting a clinical interview where details of case history could be revealed; (ii) if admitted to hospital ward doctors, nurses, etc., can observe progress/lack of it or any unusual behaviour; (iii) use of biological tests such as blood test or any other to investigate whether claim of 'illness' is confirmed or not.

- (c) Most candidates were awarded 1 mark (or 2 marks for stating two) general points about the strengths of case studies such as 'offering insight into rare disorders' and 'allowing a person to be studied in detail'. What answers tended not to do was to relate these strengths to the question about researching Munchausen syndrome. Very few candidates mentioned the case study by Aleem & Ajarim (1995) or any alternative study which would have been appropriate.

Question 7

- (a) A number of candidates scored 0 marks for this question because they provided the *conclusions* of the study rather than the *causes* of accidents as the question required. Correct answers included: insufficient supervision; poor workplace organisation; technical factors; and worker inadvertence (or worker error). Answers identifying two of these were awarded 2 marks, and any elaboration, such as percentages of each, could be awarded a further 2 marks.
- (b) Some candidates could only be awarded a partial mark because their answers were correct, but very brief. For example. Writing 'not all accidents are reported because they are only minor' would be awarded 1 mark. Candidates providing more detail, such as 'and not requiring time off work', or 'could be treated with first aid rather than hospital' were awarded the full 2 marks.
- (c) Many candidates could be awarded the full 4 marks out of 4 for their answers to this question. Often a generalisation that could be made was that there was a relatively large (or wide-ranging) sample (1 mark) and when related to the study by mentioning that there were 2964 workers (or from 4 very different industries) a second mark could be awarded. Similarly, for 'no generalisation' there was often a point followed by an example, such as that the study was only conducted in one Western country (1 mark), specifically Lodz in Poland (+1 mark).

Question 8

- (a) A few candidates did nothing more than repeat the question when stating 'profit sharing is when profit is shared amongst workers'. In order for marks to be awarded candidates needed to show some knowledge that they have learned. In this instance writing that profit-sharing is 'where workers receive a percentage of the profit' or 'a reward shared annually usually before holiday time' or 'that it is an extrinsic motivator'.
- (b) In order to be awarded full marks candidates needed to provide both parts of a closed question (the question and the answer options) in addition to providing wording of the question that asked about profit sharing. For example, writing 'Do you feel that profit-sharing gives you the motivation to work? Answer yes/no' would be awarded 2 marks. Partial answers such as a question without the answer options would only be awarded 1 mark.
- (c) In response to this question candidates could often give examples related to profit sharing without giving a strength or a weakness. For example, a candidate might write 'profit-sharing is an extrinsic motivator' which is true, but whether this is a strength or a weakness is unclear. Candidates are encouraged to think and add to the answer rather than merely state information. For example, strengths might be that profit-sharing gives workers a stronger sense of belonging to the organisation; that it can increase organisational commitment and reduce absences from work; that it can lead to increased motivation because the harder workers work the more profit they receive.

Section B

Question 9

- (a) Many candidates decided to conduct an experiment and wasted time writing about IV, DV and other aspects of experiments which were not required because the question asked for a *questionnaire* and so time should have been spent planning the details of the questionnaire. Many candidates stated nothing more than 'I would post my questionnaire online' without further elaboration. Often the questions asked were general such as 'did you like the therapy' showing no understanding of REBT at all. Candidates should ensure that, using knowledge from AS, that open questions begin with 'describe' or 'explain' and closed questions have an answer option, such as yes/no or a rating scale.

- (b)(i)** For psychological knowledge candidates often correctly referred to Ellis, although a number wrote about Beck and cognitive restructuring instead. Many candidates focused on the cause of depression and wrote in detail about the A, B and C. Whilst this was creditworthy more marks could have been obtained by writing about the therapy, Ellis's behaviour therapy and the D (disputing) and E (effects).
- (ii)** Candidates achieving full marks for this question explained how they had used the knowledge (outlined in **(b)(i)**) about Ellis's D and E to devise questions used in their online questionnaires. Some candidates were awarded 0 marks because they did not link their **part (a)** answers to **(b)(i)** and wrote about general methodological things instead.
- (c)(i)** Most candidates answered this question part correctly by providing a reason for their choice of scoring questions or interpreting the data. Those linking their answer with what they outlined in their **part (a)** plan were awarded full marks.
- (ii)** Candidates often struggled to answer this question part because they had chosen to 'add up the number of 'yes' responses' (for example) and there is no weakness in doing that. A few candidates claimed that the two people adding up the responses might have made a mistake, but this is not creditworthy as it is no more than 'counting to ten' and two 'observers' would never be used to check the reliability of addition.
- (iii)** Candidates should always include ethical guidelines in their plan. All candidates could state an ethical guideline but many could not explain why it applied to their plan or how they had implemented it in their plan. Stating; participants are given the right to withdraw; does not explain why participants are given this right in relation to rational emotive behaviour therapy.

Question 10

- (a)** This question was sometimes chosen by candidates who knew nothing about suggestive selling, or even the more generic point of purchase decisions, often resulting in Level 1 marks. Some candidates decided to ignore suggestive selling instead focusing on multiple unit pricing. Such answers were also no better than Level 1 because they did not address the question specifically. Other candidates wrote excellent answers showing very good understanding of how suggestive selling works and often linked this with detailed knowledge of the experimental method.
- (b)(i)** For psychological knowledge the most appropriate research would be that by Wansink et al. (2007). This is an 'e.g.,' study and an alternative study could be substituted. Suggestive selling is where the salesperson asks the customer if they would like to include an additional purchase or recommends a product which might suit them/their needs/usage. Some candidates wrote about multiple unit pricing and whilst this is in the same syllabus bullet point it is a different strategy from suggestive selling.
- (ii)** This question part, like all other **part (b)(ii)** questions, required an explanation to show how what was described in **10(b)(i)** informed the plan in **part (a)**. For example, if there is an outline of how suggestive selling works in **part (b)(i)** and if the procedure of suggestive selling is included in the **part (a)** plan, then this question part is simply an explanation of how the knowledge was applied in the plan.
- (c)(i)** Many candidates described the controls they had applied, such as the same person doing the suggestive selling (or not) or the same products being used or ensuring that other offers were not changed for the duration of the study. Often candidates did not explain *the reason why* they had applied these controls when this is what the question asked. Candidates should give a reason and give an example of it from their plan. Some candidates failed to apply any controls at all (absent in **part (a)**), or wrote about controls here for the first time.
- (ii)** Candidates writing about controls in this question part for the first time often applied generic weaknesses of controls, such as demand characteristics being more likely as knowledge of being in a study becomes more evident. But as no examples could be given, and as no controls were mentioned in **part (a)** only minimal marks could be awarded. It is recommended that candidates read all question parts before starting their **part (a)** answer and ensure that they address the bullet points.

- (iii) There were answers which merely stated 'I used a directional hypothesis because I predicted a direction' without giving a reason *why* a direction was predicted. Candidates awarded full marks explained that Wansink claimed that suggestive selling was successful and so this is the evidence needed for a directional (or one tailed) hypothesis to be predicted.

Question 11

- (a) Most candidates planned an appropriate way to investigate whether biochemical treatments are more effective than stimulation therapy/TENS for chronic pain, and included IV, DV, controls, and an experimental design. However, a number of answers suffered from errors such as muddling IV and DV or focusing on acute rather than chronic pain. Crucially, if a question invites candidates to consider *effectiveness*, then it is essential that their plan addresses this.
- (b)(i) For psychological knowledge the most appropriate research was to consider biological treatments and/or stimulation therapy/TENS. This was done successfully by many candidates although sometimes candidates incorrectly thought that acupuncture was a stimulation therapy.
- (ii) Candidates being awarded full marks wrote about how they used their knowledge of stimulation therapy/TENS and/or biological treatments to inform their plan. This was often done though details of a treatment programme such as the dosage of a drug and the duration of the treatment followed by participants. Top marks linked their **(b)(i)** answer with what they had done in their **part (a)** plan.
- (c)(i) An independent measures design was applied by most candidates for the reason that using repeated measure would mean both treatments being used on the same participant and one would interfere with the other confounding any result. A few candidates applied a repeated design but quickly became confused when they tried to explain why they had made this choice.
- (ii) Following from **(c)(i)** the weakness focused on an independent design with most candidates stating the generic 'there is no control over participant variables'. This is a correct answer but could only be awarded 1 mark because if an example from the plan, to show how it applies to the candidate's specific plan, is missing no further marks can be awarded.
- (iii) It is essential that ethical guidelines are included in every plan and in this instance 'ethical guidelines' was also included as a bullet point. Answers being awarded 1 mark merely stated "participants gave informed consent" without this being done in the plan. Answers being awarded 2 marks explained why participants had to give informed consent (for example) and gave the example of how they had done this in their plan, such as when they invite people to participate.

Question 12

- (a) This question was poorly answered by many candidates. There were three major flaws. First, candidates used an observation of students in a classroom rather than the named method of questionnaire. Second, many candidates used the TAT (thematic apperception test) which is a projective test involving participants explaining what is happening in ambiguous scenes. It is not a questionnaire (as the question required). Third, candidates applied a generic questionnaire about motivation rather than specifically asking about achievement, affiliation and power, which would be the only way to see which of these three was most common.
- (b)(i) Relevant psychological knowledge here was McClelland's achievement-motivation theory (1961) which suggests three work-related needs: need for achievement, need for affiliation and need for power. These three were outlined by most candidates, yet many did not mention these at all in their part (a) plan. Some candidates inappropriately described the TAT in detail.
- (ii) Many answers were awarded full marks for linking their **part (b)(i)** answer to their **part (a)** plan (which is the correct way to answer this question part). However, some candidates did not refer to the three needs at all (after outlining them in **(b)(i)**) or did not explain how the three needs informed the questionnaires they had planned in **part (a)**.
- (c)(i) Many candidates chose an opportunity sampling technique and the commonly stated reason for this choice was that teachers were available in a school. What was often lacking was *how* any teacher was actually asked to participate. For candidates choosing to use a volunteer sample, what was lacking was *how* any teacher became aware of the study to allow them to volunteer for it.

- (ii) For candidates choosing opportunity sample in **(c)(i)**, they struggled to find a suitable weakness, often reverting to a generic weakness rather than a weakness that applied to their specific study. For example, claiming that 'there will be researcher bias, where participants who 'look appropriate' are chosen' without realising that all their participants would be teachers.
- (iii) As has been mentioned, if the bullet points of the question have been included in the part (a) plan then the answer to this question is very straightforward. Many candidates were awarded full marks, but others either had not addressed types of data in their plan, or never gave more than a generic 'quantitative data can be statistically analysed' which will never earn two marks. Candidates should always include an example from their plan and elaborate beyond a basic response.